

WHITE PAPER · REPLATFORM RADAR RESEARCH

The AI Pulse Score

A reproducible, openly calibrated metric for whether AI answer engines can reach, understand, and cite a website — and whether that capability survives a platform migration.

Steven Solano · Replatform Radar

Version 1.0 · July 25, 2026

Validated on 13,838 live websites · Open dataset [CC BY 4.0]

replatformradar.com/research/ai-readiness-2026

ABSTRACT

A growing share of web discovery is mediated by AI assistants and answer engines rather than ranked search results. Whether a site appears in an AI-generated answer depends on a layer of site properties — crawler access, server-rendered content, structured data, answerable substance, and URL stability — that existing tools measure only piecemeal, and that existing AI-visibility products measure only after the fact, by observing citations. This paper specifies the **AI Pulse Score**: a 0–100 metric, computable from a single page fetch, that scores a site's AI-readiness across four objective pillars (Access, Structure, Substance, Durability). All thresholds are calibrated as percentiles of measured corpora rather than chosen by intuition, the scoring contains no model-dependent step, and the full method, dataset, and a free checker are public. We validate the metric on 13,838 live websites drawn from a random Tranco sample: it discriminates (interdecile range 55–85), and it independently reproduces four findings previously established by separate instruments. We report the first population baseline — median Pulse 71; 47% of sites grade D or F — and its distribution across 16 content-management platforms and 13 industries. The central empirical result: the web is *reachable* by AI (median Access 95) but not *machine-understandable* (median Structure 55, Substance 44).

1 Introduction

Between 2023 and 2026, asking an AI became a mainstream way to find things out. By early 2025, 34% of U.S. adults had used ChatGPT, roughly double the 2023 share^[1]; by early 2026, about half of U.S. adults used AI chatbots and six in ten read AI-generated search summaries^[2]. Gartner forecast in 2024 that conventional search-engine volume would fall 25% by 2026 as generative assistants absorb queries^[3]; whatever the exact magnitude, the direction is no longer disputed. Answer engines — ChatGPT with search, Perplexity, Google's AI Overviews, and their successors — now sit between websites and a meaningful fraction of their audience.

These systems do not present ten links; they synthesize an answer and cite a small number of sources. For a website, the operative question has therefore changed from "*where do we rank?*" to "*can the machine reach our content, parse what it means, find something worth quoting, and rely on the URL it cites?*" Each clause of that question corresponds to concrete, observable properties of a site. Yet site owners have no standard way to measure them together, and — the gap this work is specifically aimed at — no way to know whether a planned **platform migration** will silently destroy them. Replatforming is precisely the moment when structured data, server rendering, and stable URLs are most often lost, because none of them are visible to a human reviewing the new site^[16].

This paper contributes: (i) a **definition** of an AI-readiness metric designed to be objective, reproducible, and computable from one page fetch; (ii) a **calibration method** that derives every density threshold from measured percentiles of the live web rather than from judgment; (iii) a **validation** of the metric at population scale (n=13,838) including four internal-consistency checks against findings previously established with

different instruments; and (iv) the first **baseline distribution** of AI readiness across the web, by platform and by industry, published as an open dataset^[19].

2 Background and related work

2.1 How answer engines consume the web

Modern answer engines combine a language model with retrieval over web content (retrieval-augmented generation^[5]). The retrieval layer is served by crawlers that vendors document publicly — OpenAI's GPTBot, OAI-SearchBot, and ChatGPT-User^[7], Anthropic's ClaudeBot^[8], PerplexityBot^[9], and Google-Extended for Google's generative products^[10] — all of which honor the Robots Exclusion Protocol^[13]. Two measured properties of this crawler fleet motivate the metric's design. First, a joint Vercel/MERJ analysis of over 500 million GPTBot fetches found **no evidence of JavaScript execution by any major AI crawler**: content that exists only after client-side rendering is invisible to them^[6]. Second, the volume is already material — on the measured network, GPTBot generated 569M monthly requests against Googlebot's 4.5B, with ClaudeBot at 370M^[6]. In short: a large, growing crawler population reads the raw HTTP response, not the rendered page.

2.2 The machine-readable layer

The web has had a shared vocabulary for machine-readable meaning since Google, Bing, and Yahoo launched Schema.org in June 2011 (Yandex joined that November); tens of millions of domains now embed it^[11]. Answer engines inherit this layer: structured data, canonical URLs, and clean metadata are the strongest declarative signals a page can offer about what it is. On the content side, search-quality frameworks such as Google's E-E-A-T emphasize experience, expertise, authoritativeness, and trust signals — authorship, organizational identity, freshness — that are equally legible to retrieval systems^[14]. More recently, the `llms.txt` proposal (Howard, 2024) offers sites a way to hand language models a curated, markdown-native map of their content^[12]; adoption remains early but includes prominent technology companies.

2.3 Prior measurement approaches

Academic work has established that content changes measurably affect generative-engine visibility: Aggarwal et al. introduced Generative Engine Optimization and showed, over a ~10,000 query benchmark, that source-side optimizations can raise visibility in generated answers by 22–41%^[4]. A commercial category of *AI-visibility monitors* (Profound, Otterly, Peec, Scrunch, and others) measures the *outcome*: how often a brand is cited in sampled AI answers. These are complementary to, and distinct from, what is proposed here: outcome monitoring is necessarily *post hoc*, opaque to causes, and sensitive to prompt sampling. What has been missing is a *cause-side* metric — an inspectable score over the site properties that citation depends on, computable before the fact, comparable across sites, and stable enough to track through a redesign or migration. The closest precedent in spirit is Google Lighthouse: an open, deterministic, locally runnable score that became a de facto industry standard for web performance precisely because its method was public and its

results reproducible. The AI Pulse Score is designed to occupy that role for AI-readiness, with one pillar — Durability — that no existing tool measures at all.

3 Design principles

1. **Objective and reproducible.** The score contains no model-dependent step. Two parties scoring the same response bytes compute the same number. An LLM-judged content-density signal exists (§5.3) but is deliberately excluded from the score and reported alongside it.
 2. **Observable from one fetch.** Every signal is computable from a single polite HTTP request (plus two auxiliary fetches: `/robots.txt`, `/llms.txt`). This makes the metric cheap, non-intrusive, and runnable at population scale.
 3. **Calibrated, not asserted.** Wherever a signal is continuous (content depth, markup density), its scoring thresholds are percentiles of a measured corpus of the live web (§5), not round numbers chosen by the authors.
 4. **Crawler-faithful.** The metric models what AI crawlers demonstrably do: fetch raw responses without executing JavaScript^[6], honor robots.txt^[7,13], and parse declared structure^[11].
 5. **Migration-aware.** Because the moment of greatest risk to AI-readiness is a replatform^[16], the metric includes a Durability pillar scoring whether the properties that make a site citable would survive one.
-

4 The metric

The AI Pulse Score is a weighted mean of four pillar scores, each 0–100:

$$\text{Pulse} = 0.28 \cdot \text{Access} + 0.28 \cdot \text{Structure} + 0.22 \cdot \text{Substance} + 0.22 \cdot \text{Durability}$$

Letter grades follow fixed cutoffs ($A \geq 90$, $B \geq 80$, $C \geq 70$, $D \geq 55$, $F < 55$). Access and Structure carry the larger weights because they are strictly gating: content that cannot be fetched or parsed cannot be cited regardless of its quality. A fifth pillar — live citation probing against answer-engine APIs — is specified but designated *beta* and excluded from v1: it is outcome- rather than cause-side, vendor-dependent, and not reproducible from the site alone. Signal-level definitions, weights, and rubrics appear in Appendix A; the four pillars are summarized here.

4.1 Access (28%) — can AI reach the content?

Whether the twelve documented AI crawler user-agents (GPTBot, OAI-SearchBot, ChatGPT-User, ClaudeBot, Claude-User, PerplexityBot, Perplexity-User, Google-Extended, CCBot, Bytespider, Amazonbot, Applebot-Extended) are permitted by robots.txt; whether meaningful content survives *without JavaScript execution* (the extractability signal, directly motivated by the crawler-rendering evidence^[6]); HTTPS and response

health; absence of bot walls and interstitials; presence of `llms.txt`^[12] (a deliberately small, 5% weight, reflecting its early adoption).

4.2 Structure (28%) – can AI parse it?

Presence of JSON-LD structured data and coverage of high-value Schema.org types (Organization, WebSite, BreadcrumbList, Article, FAQPage, Product, Person...)^[11]; a self-referential canonical; exactly one H1; title and meta description; Open Graph tags; and the content-to-markup token ratio, percentile-scored (§5.1).

4.3 Substance (22%) – is there something worth citing?

Explicit question-and-answer structure (FAQ/QA schema or detectable Q&A blocks); freshness (most recent visible date, capped at the current year); author and organization trust signals in the E-E-A-T tradition^[14]; and content depth in tokens, percentile-scored. All four are objective checks.

4.4 Durability (22%) – will it survive a migration?

URL stability (query strings, session identifiers, mixed case, path depth — the URL features that break at cutover); whether cite-worthy pages (those carrying structured data) sit on fragile URLs; content portability, penalizing page-builder lock-in (Elementor, Divi, WPBakery, Beaver Builder, Oxygen, Bricks) whose markup does not transfer between platforms; and whether content exists as structured, reusable data rather than one-off HTML. This pillar operationalizes the observation that migrations, not redesigns of content, are where AI-readiness is most often lost^[16].

5 Calibration

5.1 Objective corpus

Continuous signals are scored against percentiles of a 3,980-site corpus fetched in July 2026 (842 random-web sites plus 1,567 higher-education and 1,571 government sites; bot-walled responses excluded). Tokenization uses OpenAI's `cl100k_base` encoding^[15]. Measured distributions: content depth $p_{10}=347$, $p_{25}=624$, $p_{50}=1,007$, $p_{75}=1,628$, $p_{90}=2,772$ tokens; content-to-markup ratio $p_{10}=0.011$, $p_{25}=0.021$, $p_{50}=0.040$, $p_{75}=0.070$, $p_{90}=0.104$; JavaScript weight (share of response tokens inside script payloads) $p_{50}=0.126$, $p_{90}=0.405$. Scoring buckets are set at these percentile boundaries (Appendix B). One design consequence is worth noting: the content-to-markup ratio is small *everywhere* on the modern web (median 4%), so any fixed-threshold rubric would have scored nearly all sites identically; percentile scoring is what makes the signal informative.

5.2 Extractability tiers

The without-JavaScript content measurement across the same corpus found 7.6% of general-web sites, 2.7% of government sites, and 0.2% of higher-education sites are effectively *JavaScript shells* (<100 extractable

tokens) — near-invisible to non-rendering crawlers^[6]. The tier boundaries (100 and 300 tokens, and the corpus p25 of 624) define the extractability rubric.

5.3 The excluded LLM signal

A fifth substance signal — semantic density, the count of distinct answerable informational units per 1,000 content tokens as judged by a language model — was calibrated on a 180-site stratified sample (p10=21, p50=34.5, p90=48). It is **deliberately excluded from the Pulse Score** and reported separately, for three reasons established during this work: (i) it is not reproducible in the strict sense (model- and prompt-dependent); (ii) it is subject to provider rate limits, which in an early bulk run caused a subset of sites to be scored by a different effective method than the rest — an inconsistency we consider disqualifying for a comparative metric and corrected by recomputation; and (iii) its exclusion makes the published score identical whether or not any AI vendor is available. On the 5,847-site subsample where it ran cleanly, median semantic density was 32.4 units per 1,000 tokens; it is published in the dataset as a supplementary field^[19].

6 Validation

6.1 Population study

The metric was computed for a random sample drawn from the Tranco top-sites list^[17] — chosen over commercial rankings for its manipulation resistance and reproducibility — yielding **13,838 measurable sites** (790 bot-walled or unreachable responses excluded rather than scored). Sub-samples of 600 and 1,406 sites scored earlier in the same week produced medians within one point of the final value, indicating the estimate converges by roughly $n \approx 1,500$ and is not an artifact of sample size.

6.2 Discriminant behavior

The score neither saturates nor collapses: interdecile range 55–85, median 71, with mass in every grade band (A 4%, B 24%, C 26%, D 37%, F 9%). Pillar medians separate sharply (Access 95, Durability 91, Structure 55, Substance 44), which is the substantive finding of §7 and also evidence that the pillars measure different things.

6.3 Internal-consistency checks

A new metric should reproduce, from its own signals, findings already established with independent instruments. Four such checks were run; all agree:

Prior finding (independent instrument)	AI Pulse result
WordPress carries ~2.6× Drupal's structured-data incidence [27k-site readiness study ^[16]]	WordPress is the top-scoring platform (median 80.9); its lead concentrates in the Structure pillar

Government websites are the least AI-citable segment [9,324-site .gov study ^[18]]	Government is the lowest-scoring sizable industry (median 66.4)
Joomla trails all major platforms on readiness signals ^[16]	Joomla is the lowest-scoring platform (median 64.8)
Institutional sectors trail commercial sectors on machine-readable structure ^[18,20]	Education, nonprofit, and government occupy three of the four lowest industry medians

Each prior finding was produced by a different measurement pipeline (signal presence counts, not Pulse scoring) on a different sample. Agreement across instruments is the strongest evidence available short of external replication — which the open dataset and free checker exist to invite.

7 Findings: the state of AI readiness, 2026

7.1 Headline

The web is reachable, but not machine-understandable. Nearly every site lets AI crawlers in and serves them substantial server-rendered content (median Access 95); URLs are largely stable (median Durability 91). But the median site scores 55 on Structure and 44 on Substance — the two pillars that determine whether an answer engine can parse a page and find something quotable in it. **46.8% of sites grade D or F overall.** The constraint on AI visibility, at population scale, is not access; it is meaning.

7.2 By platform

Platform	Sites	Median Pulse
WordPress	2,569	80.9
Salesforce Commerce Cloud	56	77.8
Shopify	24	77.5
Optimizely (Episerver)	54	76.5
Wix	23	76.4
Concrete CMS	24	76.2
HubSpot CMS	259	76.1
Adobe Experience Manager	386	76.0
Squarespace	17	73.4
Adobe Commerce (Magento)	30	72.8
TYP03	81	71.6

Sitecore	127	69.6
Drupal	492	69.1
Custom / unidentified	9,621	68.3
Kentico	26	67.1
Joomla	20	64.8

The gradient is legible: **commerce platforms and managed site builders lead**, because their templates emit product and organizational schema by default; **legacy open-source platforms trail**, where structured data is opt-in. Platform defaults, not owner diligence, appear to be the dominant determinant of a site's AI-readiness — consistent with the "defaults become blind spots" mechanism documented in the web-wide readiness study^[16].

7.3 By industry

Industry	Sites	Median Pulse
Travel & hospitality	84	75.5
Media, news & publishing	858	75.2
Real estate	20	74.7
Ecommerce & retail	349	73.9
Healthcare & medical	108	72.7
Finance & insurance	172	71.9
Business & professional services	188	71.6
Technology & SaaS	1,388	71.6
Entertainment & gaming	675	69.3
Education	377	69.1
Nonprofit & community	155	68.4
Government & public sector	193	66.4
Blog / personal	31	59.4

Industry labels were assigned by a language model from homepage title, description, and text on a random 4,598-site subsample; medians use the same objective Pulse as the rest of the study. Sectors whose economics depend on being found — travel, media, commerce — lead; institutional and personal sites trail. The inversion is pointed: the sectors publishing the most authoritative information (government, education) are the ones answer engines can least reliably parse^[18,20].

8 Limitations and threats to validity

- **Homepage-only.** All signals are measured on the homepage. Content-depth and substance signals are therefore floors: article and product pages typically carry more schema and more quotable content than navigation-oriented homepages. Multi-page sampling is the highest-value methodological extension (§9).
- **Presence, not quality.** "Has structured data" means valid-looking JSON-LD was detected, not that it is complete, accurate, or honored by any particular engine.
- **Fingerprint coverage.** 70% of sampled sites expose no platform fingerprint and are grouped as custom/unidentified; per-platform medians describe fingerprintable deployments only, and conservative detection means platform counts under- rather than over-report.
- **Crawler behavior is a moving target.** The extractability signal reflects the measured 2024–25 state of AI crawlers (no JavaScript execution)^[6]. If major AI crawlers begin rendering, its weight should fall; the versioning policy (§9) exists for exactly this.
- **Industry labels are model-assigned** on a subsample, with the error modes of zero-shot classification; the industry cut should be read as robust ordering, not precise shares.
- **Sample frame.** Tranco's top-sites frame over-represents popular, English-language, and technology-adjacent sites relative to the whole web; absolute medians should be quoted with that frame stated.
- **Authorship interest.** Replatform Radar sells migration-readiness scanning. The mitigations are structural: every threshold is derived from published corpora, the scoring code path is objective, the dataset is open, and the checker is free — any party can recompute any number in this paper.

9 Reproducibility, use, and versioning

Open artifacts. The population dataset (JSON/CSV, CC BY 4.0), the full method notes, and this paper are published at replatformradar.com/research/ai-readiness-2026. Any site's score can be computed interactively, free and without signup, at replatformradar.com/pulse; the same four objective pillars, identically weighted, produce the same number as this study's pipeline.

Versioning. The metric is versioned semantically. v1.0 fixes the four pillars, the 28/28/22/22 weights, the signal set, and the calibration constants of §5. Recalibration against fresh corpora that shifts any bucket boundary, or any change to signals or weights, increments the version, and scores are always reported with their version. Two changes are anticipated for v1.x: multi-page sampling, and a recalibrated extractability weight if crawler rendering behavior changes^[6].

Intended use. The score is designed for three uses: (i) a pre-migration baseline and post-migration regression check — the Durability pillar makes explicit what a replatform puts at risk; (ii) longitudinal tracking of a site or portfolio; (iii) population research of the kind reported in §7. It is not a ranking predictor and does

not claim to predict citation frequency in any specific engine; it measures the site-side preconditions those engines demonstrably depend on^[4,6,11].

Summary. AI-readiness is measurable, most of the web hasn't done the work, and the gap is concentrated in meaning rather than access. Because the metric is objective, calibrated, open, and free to compute, any claim in this paper can be checked by anyone with a URL.

References

1. Pew Research Center. *34% of U.S. adults have used ChatGPT, about double the share in 2023*. June 25, 2025. pewresearch.org/short-reads/2025/06/25/.
2. Pew Research Center. *Americans and AI 2026: Chatbots, Smart Devices and Views on Impact*. June 17, 2026. pewresearch.org/internet/2026/06/17/.
3. Gartner, Inc. *Gartner Predicts Search Engine Volume Will Drop 25% by 2026, Due to AI Chatbots and Other Virtual Agents*. Press release, February 19, 2024.
4. Aggarwal, P., Murahari, V., Rajpurohit, T., Kalyan, A., Narasimhan, K., Deshpande, A. *GEO: Generative Engine Optimization*. Proc. 30th ACM SIGKDD (KDD '24), Barcelona, 2024. doi:10.1145/3637528.3671900; arXiv:2311.09735.
5. Lewis, P., et al. *Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks*. NeurIPS 2020. arXiv:2005.11401.
6. Knowles, M., et al. (Vercel & MERJ). *The Rise of the AI Crawler*. Vercel Research, December 2024 (updated 2025). vercel.com/blog/the-rise-of-the-ai-crawler.
7. OpenAI. *Overview of OpenAI Crawlers* (GPTBot, OAI-SearchBot, ChatGPT-User). developers.openai.com/api/docs/bots. Accessed July 2026.
8. Anthropic. ClaudeBot crawler documentation. support.anthropic.com. Accessed July 2026.
9. Perplexity. *Perplexity crawlers*. docs.perplexity.ai. Accessed July 2026.
10. Google. *Google's crawlers and fetchers* (Google-Extended). developers.google.com/search/docs/crawling-indexing. Accessed July 2026.
11. Guha, R.V., Brickley, D., Macbeth, S. *Schema.org: Evolution of Structured Data on the Web*. ACM Queue 13(9), 2015 / CACM 59(2), 2016. Schema.org launched June 2, 2011 by Google, Bing, and Yahoo; Yandex joined November 2011.
12. Howard, J. (Answer.AI). */llms.txt — a proposal to provide information to help LLMs use websites*. September 3, 2024. llmstxt.org; answer.ai/posts/2024-09-03-llmstxt.html.
13. Koster, M., Illyes, G., Zeller, H., Sassman, L. *Robots Exclusion Protocol*. IETF RFC 9309, September 2022.
14. Google Search Central. *Creating helpful, reliable, people-first content* and Search Quality Rater Guidelines (E-E-A-T). developers.google.com/search. Accessed July 2026.
15. OpenAI. *tiktoken* (cl100k_base tokenizer). github.com/openai/tiktoken.
16. Replatform Radar. *The State of CMS Migration Readiness 2026*. July 2026. replatformradar.com/research/cms-migration-readiness-2026. (n=27,385.)
17. Le Pochat, V., Van Goethem, T., Tajalizadehkhoob, S., Korczyński, M., Joosen, W. *Tranco: A Research-Oriented Top Sites Ranking Hardened Against Manipulation*. NDSS 2019. doi:10.14722/ndss.2019.23386. tranco-list.eu.
18. Replatform Radar. *The State of Government CMS Readiness 2026*. July 2026. replatformradar.com/research/government-cms-readiness-2026. (n=9,324.)
19. Replatform Radar. *The State of AI Readiness 2026* (open dataset, CC BY 4.0). July 2026. replatformradar.com/research/ai-readiness-2026. (n=13,838.)
20. Replatform Radar. *The State of Higher-Ed CMS Readiness 2026*. July 2026. replatformradar.com/research/higher-ed-cms-readiness-2026. (n=1,796.)

Appendix A — Signals, weights, and rubrics

A.1 Access (pillar weight 28%)

Signal	Wt	Measurement & rubric
AI-crawler access	35	robots.txt parsed per agent group; score = % of the 12 documented AI agents fully permitted; a site-wide <code>Disallow: /</code> for an agent (or <code>*</code>) counts that agent as blocked.
Extractability	35	Tokens of visible text after removing <code>script/style/noscript/svg/comments</code> , no JS execution. ≥ 624 (corpus p25)=100; 300–623=70; 100–299=50; <100 (JS shell)=0.
HTTP/TLS health	15	Final status 200 over HTTPS=100; degraded=50.
Not walled	10	No bot-wall/consent-interstitial fingerprint on the primary content=100; walled=0.
llms.txt	5	Present and non-empty at <code>/llms.txt</code> =100; absent=0.

A.2 Structure (pillar weight 28%)

Signal	Wt	Measurement & rubric
Structured data present	20	≥ 1 JSON-LD block=100; none=0.
High-value schema types	25	25 points per distinct high-value @type (Organization, LocalBusiness, WebSite, BreadcrumbList, Article/NewsArticle/BlogPosting, FAQPage/QAPage, Product, Person), capped at 100.
Canonical	15	Self-referential=100; present, non-self=60; absent=0.
Heading hierarchy	15	Exactly one H1=100; multiple=60; none=20.
Title + description	10	Both=100; one=50; neither=0.
Content-to-markup ratio	10	Percentile buckets: <0.021=25; 0.021–0.040=50; 0.040–0.070=75; >0.070=100.
Open Graph	5	Present=100; absent=0.

A.3 Substance (pillar weight 22%; weights renormalized over present signals)

Signal	Wt	Measurement & rubric
Answerability	30	FAQPage/QAPage schema or detectable Q&A blocks=100; none=0.

Freshness	15	Most recent visible year ($\leq$ current year): current/previous=100; $\leq 3y=60$; older/none=20; +20 (cap 100) when date metadata [dateModified/datePublished/dateTime] present.
E-E-A-T signals	15	Author markup + Organization schema: both=100; one=50; none=0.
Content depth	20	Percentile buckets: $<347=20$; $<624=40$; $<1,007=60$; $<1,628=80$; else 100.

A.4 Durability (pillar weight 22%)

Signal	Wt	Measurement & rubric
URL stability	20	100 – 30 per fragility flag (query string; session id; mixed case; path depth >4), floor 0.
Redirect hygiene	20	Entry URL served without redirect=100; redirected=70.
AI-load-bearing fragility	25	Structured (cite-worthy) page on a fragile URL (stability <70)=30; otherwise 100.
Content portability	20	Page-builder fingerprint (Elementor, Divi, WPBakery, Beaver Builder, Oxygen, Bricks)=30; none=100.
Structured content	15	Structured data present=100; HTML-blob only=40.

The instant checker at </pulse> implements A.1–A.4 with three documented deltas: (i) the redirect-hygiene signal is omitted — the checker follows redirects and scores the final URL, with the remaining durability weights renormalized proportionally; (ii) approximate token counts (chars/4) substitute for cl100k_base; (iii) bot-wall detection is approximated by the extractable-content check rather than the scanner's interstitial fingerprints. All signal weights and rubrics otherwise match this appendix.

Appendix B – Grade cutoffs and calibration constants

Grades: $A \geq 90 \cdot B \geq 80 \cdot C \geq 70 \cdot D \geq 55 \cdot F < 55$. Objective corpus (n=3,980, July 2026): content-depth percentiles 347/624/1,007/1,628/2,772; content-to-markup 0.011/0.021/0.040/0.070/0.104; JS-weight p50 0.126, p90 0.405 (>0.405 flags render-fragility). Semantic-density corpus (n=180): 21/28/34.5/43/48 units per 1k tokens — supplementary signal only. Population study: n=13,838 scored, 790 excluded; grades 4/24/26/37/9 (A/B/C/D/F); pillar medians 95/55/44/91 (Access/Structure/Substance/Durability). All constants are fixed for metric v1.0.

© 2026 Replatform Radar LLC · Licensed CC BY 4.0 — reuse with attribution. Suggested citation: Solano, S. *The AI Pulse Score: A Reproducible, Open Metric for Website AI-Readiness* (v1.0). Replatform Radar Research, July 2026. replatformradar.com/research/ai-readiness-2026.