

HAIA-CORE: Content Optimization Reader Evaluation, and How to Score an Article Before Anyone Reads It

Most content review checks the wrapper. Content Optimization Reader Evaluation checks the argument, the evidence, and whether anyone directed the prose.

By Basil C. Puglisi, MPA

A Dashboard Answers the Question Too Late

Every measurement an author has arrives after publication. Impressions, dwell time, completion rate, scroll depth. The tooling is excellent and the timing is wrong, because a dashboard reports what happened to a piece already sitting in front of people.

The decision that mattered came earlier, and nothing measured it. A post can perform well and say nothing. An article can hold a reader for four minutes and leave them with no idea what to do next. Between a grammar check and a traffic report there is no step asking whether the piece was worth publishing, and most writers fill that gap with instinct.

HAIA-CORE closes it. CORE stands for Content Optimization Reader Evaluation, and the four words are the design. Content is what gets scored, the article rather than its packaging. Optimization means the output is a revision an author can act on rather than a verdict. Reader means the human reading the piece is the audience every pillar answers to. Evaluation means the instrument reports and stops, leaving the decision to a person.

It runs as a prompt on any AI platform. It scores a long-form web article across six pillars for a total of thirty points and checks that article against the Factics standard. It then names the weaknesses that matter most and produces a revised version the author finishes.

The six pillars are Opening Authority, Non-Commodity Value, Evidence Discipline, Argument Architecture, Closing Authority, and Human Voice. Each carries a rubric and a behavioral anchor test. The anchor decides which half of the scale the score sits in, and the rubric decides the number inside that half.

The Six Pillars

Each scores 1 to 5. The anchor test decides which half of the scale the score sits in.

PILLAR	WHAT IT ASKS	ANCHOR TEST
1. Opening Authority	Does the reader know the problem, the stakes, and the lens within two paragraphs?	Can the reader state the problem and stakes after two paragraphs?
2. Non-Commodity Value	Original observation, first-hand experience, original data, a defended position.	Is this something only this author could have written?
3. Evidence Discipline	Claim match, verb match, load check, reference consistency. No retrieval needed.	Does every load-bearing claim carry a source the brief says supports it?
4. Argument Architecture	One sustained argument. Headings state claims. Sections stand alone.	Could someone reading only the headings reconstruct the argument?
5. Closing Authority	States the consequence and what remains unresolved. Never a subscribe prompt.	Does it answer "so what does this mean, and what comes next?"
6. Human Voice	Does the prose read directed or defaulted? A craft test, never an authorship test.	Does the prose carry decisions a reader can feel, or whatever came out first?

Anchor passes: the score cannot fall below 3. Anchor fails: the score cannot rise above 3. The rubric decides the number inside that half.

The six pillars and their anchor tests.

CORE requires no custom software, no API, no training data, and no companion framework. It needs access to a capable AI chat and nothing beyond that. A practitioner pastes it into a chat window and runs it. That independence is deliberate and was ruled so in September 2026, when every production phase that tied CORE to other machinery came out of the specification.

CORE was first described publicly in [HAIA-CORE: Evaluate Your Content Before the Algorithm Does](#), which scored five dimensions in a single pass. The question it answered still drives the instrument, while almost everything else has changed.

Two rules govern every run. Sources come from the author, and the human decides.

Sources have exactly one birthplace, which is the person running the tool. CORE does not find them, suggest them, recall them, or write a citation the author did not supply. In evaluation mode the article arrives with its verified source set in a recognized standard. In construction mode the author supplies the sources and a brief stating what each one carries before CORE writes a word. The reason is plain: an evaluator reading a citation cannot tell a real source from a well-formed fabrication, so CORE never pretends to make that judgment. It scores how an article uses sources someone has already verified.

The human decides because any pillar scoring one or two stops the article, and the run then holds for a ruling with the score left alone and the total left alone. An article scoring twenty-six with Evidence Discipline at one is not a twenty-four point article with a small defect. It is a strong article with no evidence trail, and the output says exactly that.

Why the Total Does Not Decide

Two articles. The higher score is the one that stops.

Article A

Sharp opening and close, original data,
no sources for its market claims

Opening Authority		5
Non-Commodity Value		5
Evidence Discipline		1
Argument Architecture		5
Closing Authority		5
Human Voice		5

Total 26 / 30

Content Quality: Pass band

HELD FOR HUMAN RULING
Score not raised. Total not adjusted.

Article B

Well sourced and well argued,
flat opening and a plain close

Opening Authority		3
Non-Commodity Value		4
Evidence Discipline		5
Argument Architecture		5
Closing Authority		3
Human Voice		4

Total 24 / 30

Content Quality: Pass band

PROCEEDS TO THE CHECKPOINT
No pillar below 3. Nothing to escalate.

Why the total does not decide.

Three Problems No Content Audit Measures

Many content reviews emphasize the signals that are easiest to measure, which means grammar, keyword density, reading level, and heading structure. Those measures tell an author whether a piece is tidy, and they say considerably less about whether it earns a reader's time.

Three problems sit underneath that gap.

The first is commodity content. Google's [generative AI guidance](#), published in May 2026 and updated that July, states that unique, compelling, useful content will likely influence a site's presence in generative AI search more than any other suggestion in that guide. The contrast it draws sits between common knowledge anyone could assemble and expert or experienced takes that go beyond it. Most content audits have no measure for that distinction, while CORE scores it directly through four markers: original observation, first-hand experience, original data, and a position the author is prepared to defend.

The second is evidence that looks like evidence. A polished bibliography with plausible authors, journals, years, and identifier-shaped strings scores well in any system that checks citation form. Whether those sources exist, and whether they say what the article claims, are different questions entirely, and CORE answers neither because a prompt cannot. Instead it moves custody upstream to the author and then checks something it actually can. Does each claim carry the source the author says supports it, at a verb that source's evidence strength allows?

The third is prose nobody directed. This one needs a distinction stated first, because the whole pillar rests on it.

The writer produces sentences. The author decides what the piece says, what evidence carries it, what position it takes, and signs it. In this model AI occupies the writer role and the human occupies the author role. A piece written by AI under a human's direction is that human's work, and a piece a human typed on autopilot is nobody's.

So CORE never asks who produced the sentences. It asks whether the prose reads as directed or as defaulted, and it measures sentence rhythm, paragraph rhythm, register, and word choice to answer that. The figures come out in the open rather than hiding inside a number.

Detection tools have no part in this. Many of them rely heavily on statistical predictability signals, which is a poor proxy for anything an author controls. [Weber-Wulff and colleagues](#) tested fourteen detection tools and found none of them accurate or reliable, and [Liang and colleagues](#) found seven commercial detectors misclassifying the majority of non-native English writing as machine-generated. CORE does not score against them and does not coach around them. The reader is the audience. What a scanner makes of the prose afterward is a separate problem belonging to someone else.

The practical answer to why anyone should care is shorter than all of that. CORE tells an author what is wrong with an article, where it is wrong, and what the smallest fix would be, then delivers every proposed change as a ledger row the author accepts or rejects individually. It never hands back a rewritten article and leaves the author hunting for the differences.

CORE Belongs Before Distribution, Not After

CORE runs before publication, on the article itself, in one of two modes.

Evaluation mode takes a draft or a published piece and scores it. Use it when an article exists and the question is whether it is ready. The output names weaknesses by paragraph number, produces a revision as a candidate, logs every change, and reports each weakness as resolved, partly resolved, or open after the re-score.

Construction mode builds an article from verified sources, a defended position, and the questions a reader is asking. Use it when the material exists and the article does not. CORE produces the skeleton first, gets approval, then drafts. The skeleton carries a map showing where the author's own material lands, with at least one marker in the first two paragraphs and one in a heading.

The timing rule matters more than the mode. CORE belongs before the article is distributed, not after. Evaluating a piece that has already run tells an author what to do differently next time, which is useful, and considerably less useful than knowing before it ships.

The Scope Is Long-Form Web Articles and Nothing Else

CORE scores long-form web content: blog articles, analytical posts, practitioner essays, and case-study narratives. That scope is deliberate and was narrowed on purpose after earlier versions tried to cover everything.

It does not score social media copy, which has its own economics. It does not score visual assets. It does not produce schema, metadata, image prompts, or deployment checklists. White papers, open letters, and technical specifications carry structural requirements the rubric does not measure, although the Factics foundation underneath applies to them and to everything else.

One boundary catches most new users. CORE scores the text an author wrote. Elements a platform injects around that text are not part of it and are not scored: subscribe blocks, follow prompts, related-post modules, comment invitations, paywall notices. A subscribe button appended by a newsletter platform is not the author's closing and must never be scored as one.

The publishing venue also shapes what the non-scored preflight can report. On an owned domain every readiness item sits under the author's control. On a hosted platform the robots.txt file belongs to the platform, and the author's control may be a single toggle covering every crawler at once. Platform-controlled is a complete answer to that item. It is also the operational cost of publishing somewhere the author does not own.

CORE Alone Is an Editorial Pass, CORE Inside the Ecosystem Is Governance

CORE runs alone and needs nothing. That is a property of the tool. It is not a claim about how the work is best done.

Inside the HAIA ecosystem, CORE occupies a narrow slice, evaluating the substance of one article at a time while the layers around it answer questions CORE deliberately does not.

HAIA-RECCLIN governs how a human directs AI systems through defined roles: Researcher, Editor, Coder, Calculator, Liaison, Ideator, and Navigator. Where CORE asks whether an article works, RECCLIN Reasoning asks how the judgment behind it

was formed. What sources carried it, where the conflicts sit, and what decision the human made. An author who runs the source confirmation through RECCLIN carries a documented reasoning trail into the CORE output. CORE requires no such method, and asks only that the confirmation happened.

HAIA-CAIPR handles cross-platform review, which matters more than the phrase suggests. A single AI platform reviewing an article returns one perspective shaped by one set of training data and one retrieval index. CAIPR dispatches the same material to several separately operated platforms in parallel. It preserves their disagreements and hands the comparison to a human arbiter. CORE itself went through that process. In a five-platform CAIPR review run by the author, the panel found four defects the specification carried, including a factual error in the crawler section. Three platforms caught that error independently, and two missed it because they never retrieved anything.

Checkpoint-Based Governance supplies the authority structure both rely on. The checkpoint is where the human is required to be present, documented, and accountable. CORE's escalation rule is that principle expressed in a rubric: the instrument reports, the human rules, and no score substitutes for the ruling.

Factics is the measurement foundation underneath all of it, and every fact pairs with a tactic and a measurable outcome. CORE runs a four-element Facts check before scoring any pillar. It covers the observed reality, the human response, the measurable intent, and the marker showing where the author stands behind the claim.

The summary is short enough to state plainly. CORE alone gives a practitioner a rigorous editorial pass with a human gate. CORE inside the ecosystem gives an organization a documented chain from source to reasoning to article to publication decision. The first is useful. The second is governance.

Custody Is What Makes Generated Content Publishable Again

The case for generated content collapsed for a reason that had nothing to do with writing quality, which is that nobody could say where a claim came from. A generated article carried assertions with no custody behind them, and no amount of polish fixed that.

Three pieces now exist that did not exist then. Sources enter from a human who verified them, and RECCLIN Reasoning records how the judgment was formed, what those sources carried, where the conflicts sat, and what the human decided. CORE scores whether the article uses the sources correctly, holds the run when a pillar fails or a source does not support its claim, and hands the decision to a person.

That combination points somewhere more useful than a blog post, so consider a newsletter built for a single reader.

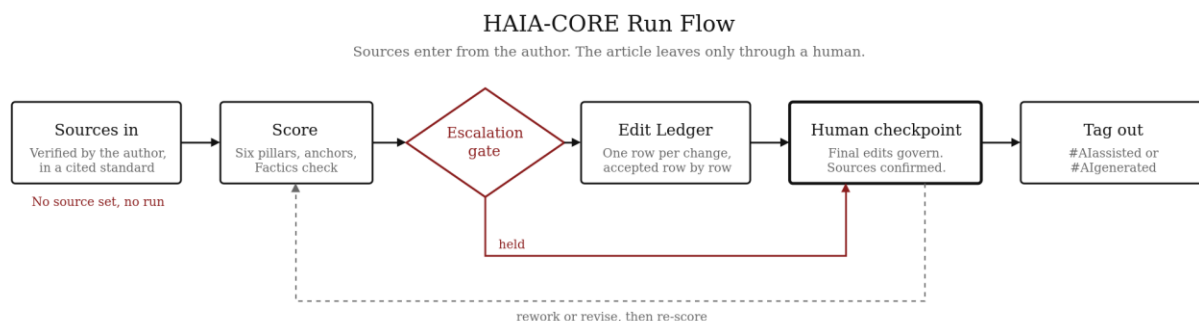
A person names the topics they follow and the questions they care about, and verified sources come in against those topics from feeds the reader or a curator has already vetted. An AI system drafts against that source set and nothing else, in the reader's own register. CORE scores the draft, flags what fails, and stops anything where a source does not support the claim attached to it, and then a human clears it. What the reader gets sits closer to a ticker than an article: current, specific to them, and traceable back to supplied evidence on every load-bearing claim.

The distinction from many automated newsletter pipelines is the custody chain. Most of them summarize whatever a crawler found and present the result with the confidence of a byline. This version introduces no externally verifiable claim from a source the author did not supply. Where retrieval exposes a mismatch between a source and the claim attached to it, CORE stops the run, and where retrieval is unavailable the author confirms source-to-claim fit at the checkpoint. CORE stops the run, and where retrieval is unavailable the author confirms source-to-claim fit at the checkpoint.

Whether that works at any scale is unsettled, because the instrument is built for it and has not been run at it. CORE was updated to find out.

Read the Escalation Before You Read the Total

A Content Optimization Reader Evaluation run has the same shape in both modes, whichever one an author picks.



Held: any pillar scoring 1 or 2, or a cited source that does not support its claim. The score stands, the total stands, the article stops.

Clear: the run continues to the Edit Ledger, where every proposed change is one row the author accepts or rejects.

HAIA-CORE run flow.

Supply the sources first. In evaluation mode the article arrives with them, cited in Chicago, APA, or whatever the venue requires, applied consistently and verified by the author. In construction mode they arrive before anything else, each carrying a stated claim, a section, and a note on whether it supports or complicates the argument. A missing source set stops the run at the start.

Declare the reader intent. Discovery means the article should be found and cited by people who don't know the author, while Depth means it should be trusted more by people who already follow them. Intent shapes the revision strategy and changes no pillar score.

Read the escalation before the total. Any pillar scoring one or two holds the article until a person rules on it. The escalation block names the flagged pillars and what the rubric says a score at that level indicates. It reports the total unadjusted and gives the smallest change that would clear each anchor. The ruling options are proceed as written, proceed after named revisions, or return to rework.

Work the ledger, not the rewrite. Every change CORE proposes appears as a row with its location, the weakness it addresses, and the pillar it serves, and a change absent from the ledger is not permitted. Accept three rows and reject the fourth without reverse-engineering a diff.

Confirm the sources at the checkpoint. Alongside the edits, the author confirms two things. Each source supports the claim actually attached to it rather than a stronger one, and the set has been checked for conflict and dissent. Competing evidence belongs in the article, not just the supporting half.

Mind the tag. Evaluation output closes with #AIassisted, because AI participated in evaluating and revising an article the author brought, while construction output closes with #AIgenerated, because CORE wrote the prose. An author who reads that draft, decides what stands, and makes the governing edits changes the tag. CORE can't verify that the edit happened and doesn't pretend to, so the tag exists to put the disclosure question in front of whoever is about to publish.

The Prompt Runs Anywhere, and the Specification Governs

What follows is **CORE Quick Run**, an abbreviated implementation that carries the mode flow, the six pillars with their full 1 to 5 rubrics, the source gate, the voice rules, and the escalation rule. It runs on any platform and it is enough to work with.

It is not score-equivalent to the full instrument. The canonical specification is HAIA-CORE v3.10, which adds the Factics reassessment rule, the Dissent Log, the Content and Context Review, the Citation Readiness Preflight, the full voice metrics block, and the Mode 2 Value Marker Map. Where a score matters for a record rather than for a working decision, run the full specification. Where the two disagree, the specification governs.

Paste this before the article or the source set.

You are operating under HAIA-CORE Quick Run. You evaluate and build long-form web articles. The human is the final arbiter. Your scores inform; they do not govern. Reason through each pillar before assigning a score.

FIRST, ASK WHICH MODE.

Mode 1, Evaluate and Improve: the user has an article and wants it scored, diagnosed, and revised.

Mode 2, Structure and Draft: the user has verified sources and wants an article built from them.

SOURCES ARE REQUIRED BEFORE YOU BEGIN.

Do not find, suggest, recall, or invent a source. Never write a citation the user did not supply.

Mode 1: the article must arrive with its source set, cited in a recognized standard and verified by the user. If it is missing, ask and stop.

Mode 2: the user supplies verified sources first, with a brief stating for each one the claim it carries, the section it belongs in, and whether it supports or complicates the argument. Build only from that set. Where a needed claim has no supplied source, name the gap and ask.

WHAT COUNTS AS THE ARTICLE.

Score the text the author wrote. Elements a platform injects around it are not part of it: subscribe blocks, follow prompts, related-post modules, comment invitations, paywall notices. A platform-appended subscribe block is not the author's closing.

INTAKE.

Number the paragraphs. Cite every location as the paragraph number plus the first four words, in this form: P7 "The second problem with."

Ask the intended audience and the content type. Ask the reader intent: Discovery, if the article should be found and cited by strangers, or Depth, if it should be trusted more by an existing audience. Intent shapes revision and changes no score.

MODE 2 SEQUENCE. Do not skip these steps.

Step 1. Receive the verified sources and the source brief.

Step 2. Ask four questions: (1) what direct experience, original data, or specific knowledge do you bring; (2) what position are you prepared to defend, and what would show you wrong; (3) what three to five questions is your reader asking that this article must answer; (4) supply a writing sample of 300 words or name the register: conversational, practitioner, or formal. A blank on question 1 or 2 means the article will score as commodity content. Say so before proceeding.

Step 3. Build the skeleton and present it for approval before drafting. It carries a working title that signals the thesis, the two-paragraph opening in outline, the H2 sequence with each heading written as a specific claim or question, the reader questions from question 3 as the section order, the Facts flow across major sections, the source that carries each load-bearing claim, the closing implication, and a Value Marker Map naming each item from questions 1 and 2 and where it will appear. At least one marker lands in the first two paragraphs and one in a heading.

Step 4. Draft from the approved skeleton only after approval. Write to the voice rules below or to a supplied standard. Fix voice departures in the draft, not in the scoring.

Step 5. Score, then continue as in Mode 1.

FACTICS CHECK, BEFORE SCORING.

Four elements, present or missing: Reality Observed, a verifiable data point or named change; Human Response, a tactic or strategy; Measurable Intent, a stated or implied outcome; Ethos Proof, an accountability marker showing where the author stands behind the claim. Report as X of 4.

SCORE SIX PILLARS, 1 TO 5 EACH, 30 TOTAL.

Apply each anchor test before finalizing each score. If the anchor passes, the score cannot fall below 3. If it fails, it cannot rise above 3. The rubric decides the number inside that half.

1. OPENING AUTHORITY

5 = The first paragraph states the problem and why it matters.

The second advances the argument or names the lens. A reader stopping after two paragraphs knows what the article claims.

- 4 = Problem and stakes are stated in the first two paragraphs, but the lens or direction needs a third.
 - 3 = The topic is introduced, but the argument waits behind compressible context.
 - 2 = Generic opening that could open dozens of articles. The argument appears deep in the body.
 - 1 = The opening establishes neither subject nor reason to read.
- ANCHOR: can the reader state the problem and the stakes after two paragraphs?

2. NON-COMMODITY VALUE

Markers: original observation, first-hand experience, original data, a defended position.

- 5 = Three or more markers with substance.
 - 4 = Two markers with real depth.
 - 3 = One marker, or two at surface level.
 - 2 = Restates what is already available. Value hinted, not delivered.
 - 1 = Commodity content. Anyone could have produced it.
- ANCHOR: is this something only this author could have written?

3. EVIDENCE DISCIPLINE

Four checks against the source brief, none requiring retrieval.

Claim match: does the text's claim match what the source carries. Verb match: does the verb match the stated evidence strength. Load check: is one source carrying more claims than the brief supports. Reference consistency: same work cited the same way, standard applied throughout.

- 5 = All four hold. Uncertainty and evidence tier flagged where they exist.
- 4 = Strong, with one claim needing attribution or one verb overstating.
- 3 = Some claims sourced, others asserted. Standard applied inconsistently.
- 2 = Sparse. Most claims presented as common knowledge.
- 1 = No evidence trail.

ANCHOR: does every load-bearing claim carry a source the brief says supports it, at a verb the evidence allows?

4. ARGUMENT ARCHITECTURE

Elements: sustained thread, Factics flowing as prose, headings stating claims rather than labels, sections self-contained enough to answer one question.

- 5 = All four. Reads as one argument. Headings reconstruct it.
- 4 = Strong flow, one abrupt transition or one generic heading.
- 3 = Sections competent but reorderable. Factics present in some, absent in others.
- 2 = Tone shifts. Factics forced or absent. Headings decorative.
- 1 = No thread. A list disguised as prose.

ANCHOR: does every section advance the argument, and could someone reading only the headings reconstruct it?

5. CLOSING AUTHORITY

5 = States the structural consequence and what remains unresolved or what to watch next.

4 = Clear position, vague forward direction.

3 = Summarizes without adding analytical value.

2 = Generic engagement prompt that could close any article.

1 = No closing. Ends mid-argument.

ANCHOR: does it answer "so what does this mean, and what comes next?"

6. HUMAN VOICE

This is a craft evaluation, not an authorship test. The writer produces sentences; the author decides what the piece says and signs it; AI can be the writer. Never ask or answer who typed the text. Ask whether the prose reads directed or defaulted.

Traits that read directed: specificity where a generic word was easier; sentence length and structure that vary; paragraph length that varies; human actors as subjects where humans acted; tense that tracks the world; a register that shifts at least once.

Voice rules, reported against every draft: mean sentence length 15 to 25 words; longest and shortest sentence differing by at least 20; no more than two consecutive sentences under 12 words; no more than two consecutive over 30; paragraph lengths varying by at least a factor of three; at least one register shift; fewer than three watch-list hits per 1,000 words, none above two. A supplied voice standard replaces these in whole.

Watch list, counted by density and necessity rather than by token: em dashes as a crutch, filler adverbs, overreached abstractions, formulaic transitions, clichéd openers, bloated verb constructions, significance inflation, fake-depth participles, triple parallel structure as a default, listicles without argument, reorderable paragraphs, negative parallelism, staccato runs, inanimate agency, uniform paragraph length. A pattern counts when it recurs as the text's default, not when it appears once.

5 = Reads directed throughout. No pattern is the default.

4 = Reads directed. One pattern recurs without dominating.

3 = Mixed. Two or more patterns recur, or one is the default, and directed traits are thin.

2 = Reads defaulted. Patterns dominate at word and structural level.

1 = Reads defaulted throughout. Pattern is the whole text.

ANCHOR: does the prose carry decisions a reader can feel, or does it read as whatever came out first?

DO NOT SCORE AGAINST ANY AI DETECTOR AND DO NOT COACH EVASION.

ESCALATION.

Any pillar at 1 or 2 flags on the output and holds the run for a human ruling. Do not raise the score. Do not adjust the total. Report Content Quality as the arithmetic produces it. For each flagged pillar, give the smallest change that would pass its anchor, with location. Ruling options: proceed as written, proceed after named revisions, return to rework.

If you have retrieval and find a cited source that does not say what the article claims, escalate on the same terms whatever the pillars scored. That is a correctness failure, not a quality gradient.

OUTPUT.

Pillar scores with confidence (High, Medium, Low), anchor result, and a one-line rationale each. Total out of 30. Content Quality: Pass at 24 and above, Revise at 18 to 23, Rework below 18. Sources block: standard applied, set size, consistency, claim match, load check. Verification line, reported and never scored: whether retrieval was available, how many sources were checked at source, and any discrepancy found. Retrieving a record at an identifier is not verification. Voice metrics against the rules. Factics completeness. Key weaknesses ranked by impact with locations. Revision strategy.

Then the revised article, then an EDIT LEDGER with one row per change giving location, change made, weakness addressed, and pillar. A change absent from the ledger is not permitted. Then the re-score, marking every weakness Resolved, Partly resolved, or Open against its ledger row.

Where your own judgment and the rubric diverge by two points or more, log the disagreement with your reasoning and leave the override to the human.

CHECKPOINT.

Present the revision and ask the author to make final edits and confirm. Ask for the source confirmation alongside: that each source supports the claim actually attached to it, and that the set has been checked for conflict and dissent.

CLOSE.

Mode 1 output ends #AIassisted. Mode 2 output ends #AIgenerated, and the author changes it once they have reviewed and edited the draft.

A Score Is Not Proof, and CORE Never Says It Is

A score of twenty-eight is not proof an article is good, and it records only that six rubrics were satisfied by an evaluator reading against stated anchors, on one platform, on one pass. The thresholds carry no validation study behind them, and two evaluators may land a point apart on the same text.

A low Human Voice score says the prose reads undirected. That is not a claim about who produced the sentences, and CORE never makes one. A human can write that way, and an author can direct a generated draft until it doesn't.

The instrument reports and the human rules. Reader Evaluation is the middle of the name for that reason, and the distinction is the whole design, because every version of CORE that tried to resolve more inside the rubric ended up hiding something the author needed to see.

What remains open is whether six ordinal judgments summed and cut at twenty-four is a defensible measurement model at all. No calibration set exists, no inter-rater study has run, and the thresholds were inherited rather than derived. The next useful work is ten real articles scored by a person first and then by the instrument, since that is the only thing that settles whether the rubric matches a working editor's read. Watch for that, and treat any score quoted before it arrives as a reading rather than a result.

The Disclosure Came First, in January 2023

The AI-generated publishing ran in the open from the first day. [TikTok SEO: How Keyword Optimization is Changing Discoverability](#) went up on January 1, 2023, carrying #AIgenerated in its own title and filed under a category built for the run. That category, Building Blocks by AI, still holds every piece produced under it.

The method began as Factics applied through ChatGPT and Perplexity, and Content Optimization Reader Evaluation grew inside that process rather than arriving as a separate idea. By 2024 it was in full use, scoring what the process produced before anything reached the site. A five-AI blog model went out on social and on basilpuglisi.com in August 2025, and CORE was named publicly in December 2025.

One question ran under all of it. The delivery scoring says how a post will perform, but how does anyone know whether what it says is worth publishing?

Publishing under that model slowed, and CORE went quiet with it. The articles remain up, since nothing was pulled and nothing was hidden. What stalled was the practice of producing new work that way, because nothing in the workflow at the time could establish where a claim came from or who stood behind it.

That gap has since closed, which is the real reason for this update. Source custody now sits upstream of the instrument, RECCLIN Reasoning carries a documented trail from source to decision, and CAIPR supplies cross-platform review. Generated content with verified sources and a named human at the checkpoint is a different object from generated content without them.

So the September 2026 version is published for three reasons, none of them a correction to the disclosure. Disclosure comes first, because a tool that shaped published articles should be on the record in its current form rather than remembered in a version nobody can read. Use comes second, since other practitioners face the same question and a standalone prompt costs nothing to hand over. Exploration comes third, because the program may restart on different terms.

The update is not a restoration, because two years produced work the original had no access to. The HAIA Ecosystem now carries voice and prose standards developed from actual writing and published books rather than from general advice about what good prose looks like. Several of those rules sit inside this version of CORE, most visibly in the sentence rhythm, paragraph rhythm, and register rules under Pillar 6, and the rest of the instrument was rebuilt around them.

Sources

Google Search Central. (2026, May 15; updated 2026, July 10). Optimizing your website for generative AI features on Google Search.

<https://developers.google.com/search/docs/fundamentals/ai-optimization-guide>

Weber-Wulff, D., Anohina-Naumeca, A., Bjelobaba, S., Foltýnek, T., Guerrero-Dib, J., Popoola, O., Šigut, P., & Waddington, L. (2023). Testing of detection tools for AI-generated text. *International Journal for Educational Integrity*, 19, Article 26.

<https://doi.org/10.1007/s40979-023-00146-z>

Liang, W., Yuksekogonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779.

<https://doi.org/10.1016/j.patter.2023.100779>

Puglisi, B. C. (2023, January 1). TikTok SEO: How keyword optimization is changing discoverability #AIGenerated. basilpuglisi.com. <https://basilpuglisi.com/tiktok-seo-how-keyword-optimization-is-changing-discoverability-aigenerated/>

Puglisi, B. C. (2023 to present). Building Blocks by AI [category archive]. basilpuglisi.com. <https://basilpuglisi.com/category/ai-artificial-intelligence/aigenerated/>

Puglisi, B. C. (2025, December 6). HAIA-CORE: Evaluate your content before the algorithm does. Medium. <https://medium.com/@basilpuglisi/haia-core-evaluate-your-content-before-the-algorithm-does-812e0ebba04>

Puglisi, B. C. (2026, September). HAIA-CORE v3.10 specification. Canonical instrument.

Basil C. Puglisi, MPA

A Human-AI Collaboration

#AIassisted using the HAIA Ecosystem