



AI assembles the raw material. People make the accountable evidentiary judgments.

2

user experiences

5

pipeline stages

11

reliability providers

–100 / +100

Verimeter scale

Describes the product as built in the alpha as of August 2026. Companion to the Investor Plan.

01 / OVERVIEW

What VeriStrata Makes

One artifact: a network of statements connected by attributable human judgments.

Most tools hand you a verdict. VeriStrata hands you the reasoning and lets you inspect it.

The platform turns an article, post, video, or podcast into an evidence map: the assertions it makes, the sources bearing on those assertions, the statements inside those sources, and the human judgments connecting them — supports, refutes, contextualizes, or insufficient evidence.

The words this document uses

Term	What it means
Assertion	A single checkable statement extracted from a piece of content
Claim	An assertion that is being tested. The same statement can be a claim in one map and evidence in another
Link	A human judgment connecting one assertion to a claim, together with how strongly
Evidence map	Every link bearing on a single claim
Evidence graph	All the maps together. The corpus, and the durable asset
Verimeter	A score from –100 to +100 computed from an evidence map
Versed	Content that has been mapped and carries a Verimeter

02 / THE READER

Five Seconds, In the Page

The reader does not come to VeriStrata. VeriStrata arrives where they are already reading.

A reader lands on a claim they are not sure how seriously to take — a piece about blue light and glasses that supposedly protect sleep, a public-health claim, a political allegation, a study, a viral thread. Today they have two bad options: trust it, or leave the page and start their own investigation. VeriStrata creates a third.

If the content has been versed, the browser extension surfaces the evidence signal without the reader leaving the page: the current Verimeter, the assertions evaluated, and the evidence on both sides of each. They can skim, inspect a source, decide the article is not worth reading, find a better one, or open the full workspace.

One pass, concretely. The article argues that blue-blocking glasses meaningfully improve sleep. The extension returns a Verimeter of –45. Five assertions have been evaluated and the central one sits at –62 — four supporting assertions from two small lab studies, against nine refuting or qualifying assertions from seven sources, including a meta-analysis the article never cites. The reader has not been told what to think. They have been shown where the evidence sits.



VeriStrata Product Introduction

The Evidence Layer | How the platform works



The same mechanism on a different article: a 5G health piece at -29%, each assertion carrying its own reading and the source that refuted it named beneath.

VeriStrata does not ask the reader to trust the answer; it shows them the argument. The product is not the rating. The product is the reasoning trail.

If the content has not been versed, the reader can submit it. The pipeline in section 04 prepares it for review.

03 / THE EVALUATOR

Drawing the Lines of Reasoning

The evaluator does not start from a blank page. They start from assembled material.

The second user is working on purpose: a student with an assignment, a researcher, a journalist, a trained community reviewer. They enter a workspace where the raw material is already organized — the assertions, the retrieved sources, the source assertions, and proposed relationships between them.

Their task is not to fact-check the internet. It is to judge which assertions bear on which claims, and how. They decide whether a source assertion supports, refutes, contextualizes, or fails to bear on the claim, and record rationale and confidence as they go. Other evaluators inspect, dispute, improve, or reuse any of those judgments — which is what turns one person's reading into a reviewed one.



An evidence map in the workspace: four claims, six source assertions, and the weighted links an evaluator drew between them, with the author and publisher of every source attached.

What emerges is durable. The evidence graph records how a claim was evaluated, what evidence was used, who made which judgment, and what the next reviewer can inspect, challenge, improve, or reuse.

04 / THE PIPELINE

What Happens Before a Person Sees Anything

Five stages prepare the evidence. Every output stays a proposal until someone judges it.

Building an evidence graph by hand is prohibitively slow, which is why nothing like it exists at scale. The pipeline that populates it automatically is the hardest thing VeriStrata has built, and each stage is running in the alpha today.

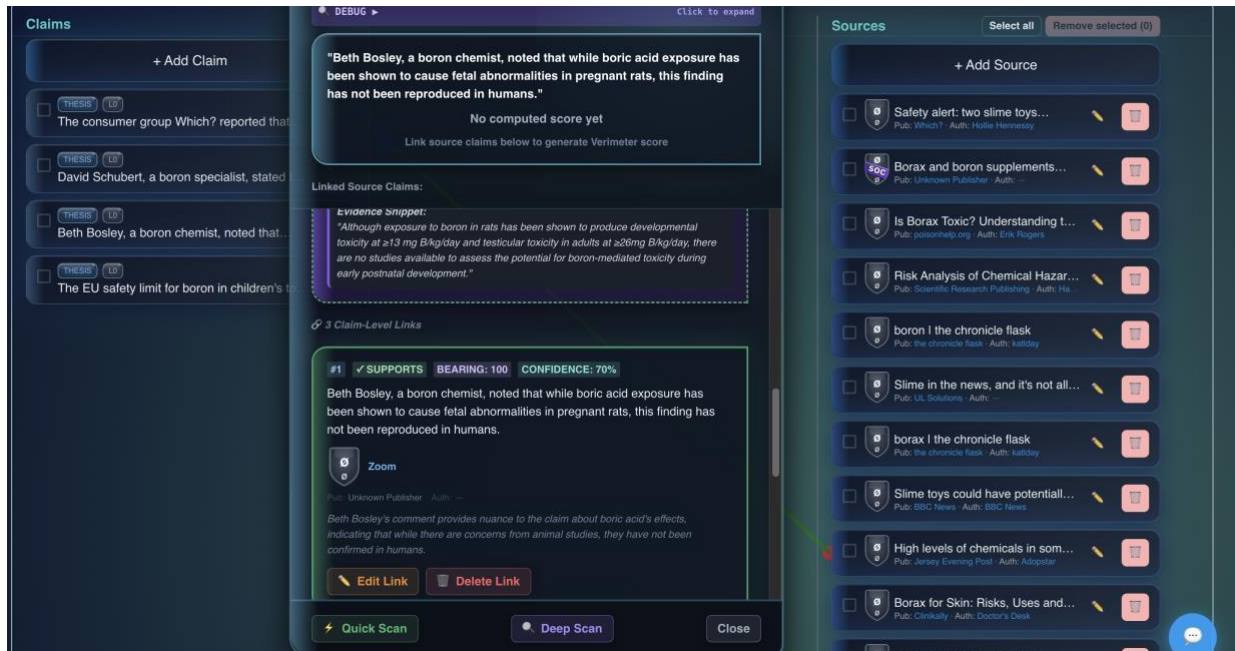
Stage	What happens	Why it is difficult
Assertion extraction	Content is decomposed into the checkable statements the author actually made	Summarizing is not extracting. Paraphrase drift produces a claim easier to refute than the one the author made, and everything built on it is worthless
Screening	Candidates are filtered before full processing	Without it, cost scales as claims × sources × assertions. This is what holds cost per map inside the margin ceiling
Evidence retrieval	Relevant sources are located across scientific literature, publications, and the open web	Search optimizes for relevance to a query. This must optimize for bearing on a proposition, and deliberately surface refuting evidence the content omitted
Source decomposition	Each source is broken into its own assertions	A study is a dozen assertions and only some are germane. Too coarse and nothing connects; too fine and the graph fills with noise
Alignment	Each source assertion is matched to each claim with a proposed link and confidence	The hardest stage. One finding can support a weak version of a claim and refute a strong one, and the system must not collapse that distinction

Every output is a proposal, not a conclusion. Evaluators accept, reject, re-scope, or reverse any of it. Model error degrades into a rejected suggestion rather than a published verdict, which is what makes the system safe to put in a classroom and defensible when it eventually rates a publisher.



VeriStrata Product Introduction

The Evidence Layer | How the platform works



A single link under review. The system proposed the relationship, its bearing, and its confidence; the evaluator accepts, re-scopes, or deletes it.

05 / THE VERIMETER

How the Score Is Computed

The score is computed from the links, not authored over them.

The Verimeter runs from -100 to +100, shown in the product with a percent suffix. At -100 the mapped evidence strongly refutes the claim; at +100 it strongly supports it. Zero means the evidence is balanced, or too thin to lean either way — and which of those it is stays visible in the map rather than hidden inside the number.

Nobody at VeriStrata decides what an article scores. Each link carries a direction and a strength, and is then weighted by two independent factors.

Source reliability

Reliability is not VeriStrata's opinion of who is trustworthy. It is assembled from independent providers, each covering a different kind of source.

Provider	What it contributes
Wikidata	Entity identity — who a publisher or author actually is
Wikipedia Perennial Sources	Community-maintained reliability assessments
SCImago	Journal impact ranking for academic sources
Crossref / OpenAlex	Scholarly metadata, DOI resolution, citation context
Media Bias/Fact Check	Independent media reliability and bias assessment
Global Disinformation Index	Risk assessment for news domains
GDELT	Global coverage and event context
Internet Archive	Archival footprint. No reliability signal, but a thin footprint caps the grade an obscure source can reach
RDAP	Domain registration provenance
OpenSanctions	Sanctions and enforcement exposure
BBB, CFPB, court records	Complaint, enforcement, and litigation history for commercial publishers

Coverage is uneven and the system reports that rather than hiding it. A peer-reviewed article resolves cleanly through Crossref, OpenAlex, and SCImago. A mainstream news domain resolves through Media Bias/Fact Check and the Global Disinformation Index. A government archive page may match almost nothing: in a live run against an archived CDC document, six providers returned no



match and the only signal was a weak archival footprint, which caps the grade rather than raising it. Where providers do not resolve a source, the link carries less weight and the gap stays visible instead of being filled with a default.

Evaluator reputation

The person who drew the link carries standing in the platform, earned the way it is on Stack Overflow or Yelp — through contributions, completed evaluations, and whether their judgments survive peer review. An established evaluator's link weighs more than a new account's.

Levels, and movement

Each claim scores from the links bearing on it; a piece of content scores from the claims extracted from it. The model ranks claims within content — thesis, theme, pillar, evidence — each level carrying roughly ten percent less weight than the one above. The hierarchy is built into the data model; the current prototype weights all claims equally when scoring, so content-level figures today are an unweighted average.

The meter recalculates every time it is displayed, and previous readings are preserved in an append-only, tamper-evident record. Any figure quoted here is a snapshot; the map behind it keeps moving. A score that never moved would mean nobody was working.

Two readings, never blended

Measure	What it represents	What can move it
Verimeter	The evidence reading, -100 to +100	Reviewed evidentiary links, source signals, and evaluator standing
Tally	The crowd's expressed position, reported as proportions — 71 of 121 supported, 21 of 121 refuted	Moderated participant responses

The two disagree often, and that disagreement is information. A coordinated group can move the Tally, because the Tally is a count of what people said. It cannot move the Verimeter without producing attributable links that survive peer review from an account with earned standing.

06 / STATUS

What Is Built, and What Is Not

Trust primitives are built and functioning at roughly thirty guided testers. None has been tested at pilot scale or against an adversary.

Governance is part of the product because VeriStrata structures how public conclusions get built. A reasonable objection is why it was built before demand was proven. The answer is that trust is not retrofittable: a system that has accumulated a year of evidence maps without reputation scoring, review thresholds, or tamper-evident records cannot acquire them later without invalidating everything already in it.

Trust primitive	Current status	Pilot validation
Tamper-evident history	Built, not load-tested	Verify concurrent classroom use and contribution history
Approval workflow	Built, not pilot-tested	Measure review consistency and disputed-map handling
Reviewer eligibility / reputation	Built, not pilot-tested	Test whether weighting improves reliability without entrenching incumbents
Anti-brigading controls	Partial / designed	Adversarial testing remains future work
Evidence corpus moat	Validation target	Depends on repeat use, quality, reuse, and scale

Why the work compounds

Every evaluator session produces a second asset alongside the evidence map: a labelled record of where a person agreed with the pipeline and where they overrode it — on assertion boundaries, on whether a source bears on a claim, on whether a link is support or context.

That is training signal for exactly the task the pipeline performs, generated as a by-product of coursework an institution is paying for. Both assets are only worth something because every judgment is attributable, which is why the governance layer was built before the demand that justifies it. It is a validation target rather than a proven moat, and it requires the pilots to happen.