

VERITROOPER Scout — System Card — SEC 10K multi — VERITROOPER v10



SEC 10K multi — VERITROOPER v10
26 August 2026

Deployed-system documentation for the evaluated AI system, generated from an independent accuracy & robustness audit. Generated 2026-08-26 17:30:00.

About this document. A system card documents the system as deployed — the model together with its retrieval and the evaluation wrapper — not just the trained weights. VERITROOPER auto-populates the measurable sections from one audit run; sections that only the provider can supply (intended use, deployment context) are marked as such and left for the provider to complete. This card is evidence of measured behaviour; it does not certify the system.

Contents (click a section to jump)

1. System identification
2. Intended use & out-of-scope uses
3. Evaluation data
4. Evaluation methodology
5. Quantitative results
6. Robustness & adversarial behaviour
7. Known limitations & failure modes
8. Remediation
9. Human oversight & sign-off
10. Reproducibility & provenance
11. Data handling & isolation posture
12. Framework alignment
13. Deployment context & risk considerations
14. Card version

1. System identification

Field	Value
Model under test	gpt-5.6-sol
Evaluated material / knowledge base	SEC 10K multi — VERITROOPER v10
Task type	generation
Evaluation engine	VERITROOPER v10
Test items	889
Card generated	2026-08-26 17:30:00

2. Intended use & out-of-scope uses

Provider-supplied — VERITROOPER does not infer intended use. Complete before publishing this card.

- Primary intended use(s): _____
- Intended users: _____
- Out-of-scope / prohibited uses: _____

3. Evaluation data

The system was evaluated on a purpose-generated question set derived from the provider's material. Question categories exercised: Exception, Negative (hallucination traps), Precision, Cause & Effect, Cross-Reference, Conditional, Calculation.

Field	Value
Test items	889
Categories	7
Refusal / unanswerable items (hallucination traps)	200
Source material	data/SEC_10K_multi.txt

Full provenance — generation method, parameters, and the exact system prompt — is in the accompanying Data Card.

4. Evaluation methodology

- **Baseline arm** — the model with standard retrieval (vanilla RAG).
- **Pipeline arm** — the same model evaluated through the VERITROOPER engine.
- **Adjudication** — each verdict checked by a verification model (the Doctor: gpt-5.6-sol) and independently audited by gemini (Gemini), applied to **both** arms under the same standard so neither arm gets a free pass.
- **Bad-test exclusion** — items whose own ground truth was malformed are excluded symmetrically from both arms; the excluded count is reported.

Independence of the judging steps

Who did what in this audit, and how independent each judging step is from the model under test (MUT):

Step	Performed by	Type	Independent of the MUT?
Answer under test	gpt-5.6-sol	the model being audited	—
First-pass check (Validator)	deterministic rule check	code, not a model	yes — not a model
Diagnosis (diagnostic reviewer)	gpt-5.6-sol	verification LLM	no — the same model as the MUT
Independent verification (Verifier)	gemini (Gemini)	model / vendor	yes — an independent verifier

Note: on this run the Doctor is the same model as the MUT, so the Doctor step is not an independent second vendor. The independent cross-check is the Verifier, applied to both arms under the same standard.

Baseline configuration

Baseline retrieval configuration — a **standardized BM25 reference baseline** the pipeline is measured against (a reproducible lexical-retrieval reference, not necessarily the customer's production RAG stack):

Parameter	Value
Retrieval	BM25 lexical (Robertson et al. defaults)
Chunking	fixed-size, 500 tokens, 50-token overlap
Retrieved chunks	top-5 by BM25 score
Reranking	none
Query rewriting	none
Dense embeddings	none

Full machine-readable run configuration is in config_snapshot.json.

5. Quantitative results

Headline Accuracy (bad-test cases excluded from denominator)

Scoring set: **889 questions** (of 889 generated; 0 identified as bad tests — questions where the ground truth itself was malformed — excluded symmetrically from both arms' denominators per audit-grade methodology.).

Metric	Without VERITROOPER (vanilla-RAG baseline)	With VERITROOPER (pipeline)
Final verified accuracy	72.89% (648 / 889)	88.08% (783 / 889)
Answers scored incorrect	241	106
Δ vs baseline		+15.19pp
Failure-Recovery Rate (pipeline recovers what % of baseline failures)		64.7%
95% confidence interval (Wilson)	69.9–75.7%	85.8–90.0%

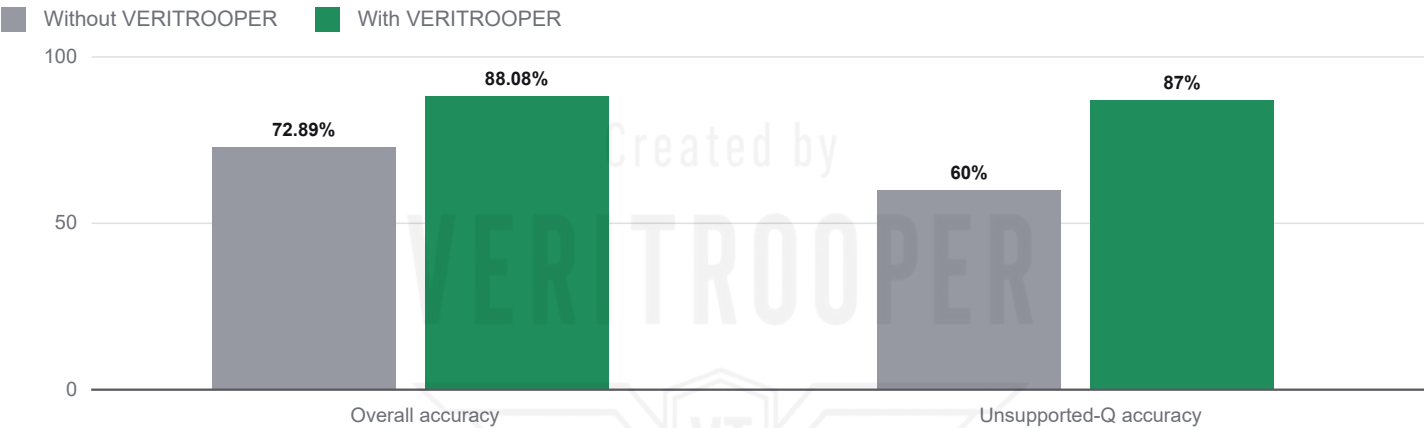
The 95% confidence interval reflects sampling uncertainty on 889 scored items: a point estimate of 100% does not prove a 0% error rate in the field — it means no error was observed in this sample; a tighter bound needs more questions.

Per-Category Accuracy & Failure Recovery

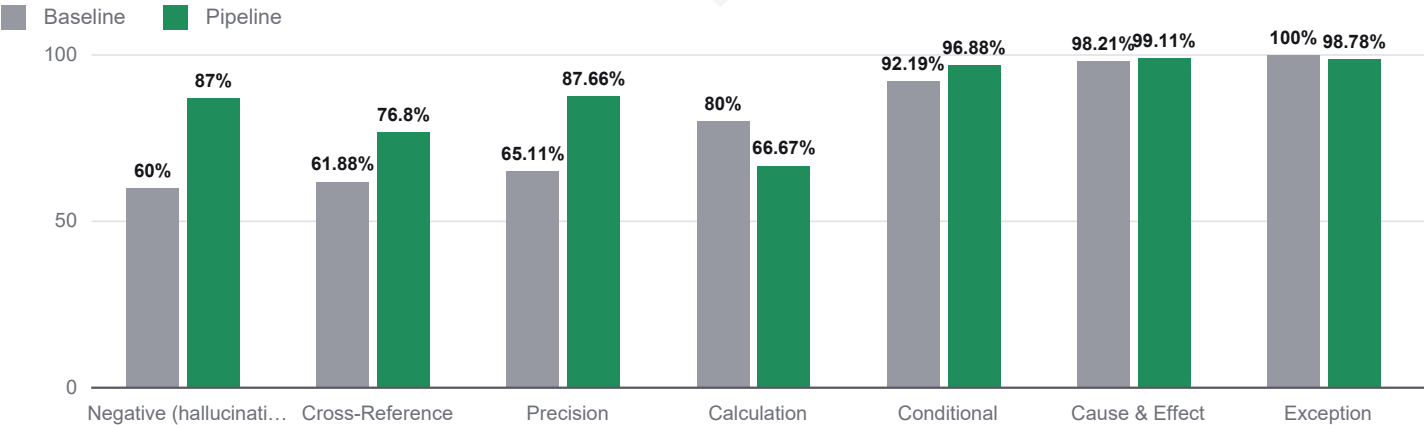
Per-category accuracy with and without VERITROOPER. **Failure-Recovery Rate** measures the fraction of vanilla-RAG baseline failures the pipeline recovers — the most direct signal of where the architecture adds value on this material. Every category is shown, including any where the pipeline matches or underperforms baseline.

Question Type	Questions	Baseline Accuracy	Pipeline Accuracy	Improvement	Failure-Recovery Rate
Negative (hallucination traps)	200	60.00%	87.00%	+27.00pp	76.2%
Precision	235	65.11%	87.66%	+22.55pp	72.0%
Cross-Reference	181	61.88%	76.80%	+14.92pp	46.4%
Conditional	64	92.19%	96.88%	+4.69pp	60.0%
Cause & Effect	112	98.21%	99.11%	+0.89pp	50.0%
Exception	82	100.00%	98.78%	-1.22pp	—
Calculation	15	80.00%	66.67%	-13.33pp	0.0%
All Categories	889	72.89%	88.08%	+15.19pp	64.7%

Accuracy — with vs. without VERITROOPER



Per-category accuracy — baseline vs. pipeline



6. Robustness & adversarial behaviour

The set includes adversarial hallucination-trap items (the negative category, n=200) whose correct response is to refuse. On these, the raw model scored **60.00%** and the pipeline **87.00%** — the single clearest signal of whether the system answers questions its evidence cannot support.

7. Known limitations & failure modes

Failure modes across the 194 confirmed baseline model errors (a failed answer may carry more than one diagnostic label — or one that maps to no named pattern — so these counts need not equal the number of failed answers):

- other — 58 answers
- deterministic:answered unanswerable — 52 answers
- deterministic:gt number conflict — 18 answers
- knowledge gap — 18 answers
- misinterpretation — 13 answers
- table misread — 11 answers

Results establish measured behaviour on the tested material, not field performance outside the tested distribution.

8. Remediation

Deterministic guardrails

Code-level checks you can deploy in your own stack to catch the failure patterns found in this run — no model retraining required. Because they are deterministic, they behave identically on every run. **Block** = the check is authoritative and can reject or correct the answer automatically; **Flag** = the check detects the error class and routes it to human review or abstention. This complements the remediation recommendations above — it does not replace them.

Block — reject or correct automatically:

- **Unit/format normalization** — Parse the value AND its unit, convert to a canonical unit/scale (% , \$, thousands vs millions) before comparison, and reject unit or scale mismatches. (Addresses: Misinterpreted, Out of scope, Table misread, Unit/format mismatch, Wrong calc/quote, Wrong period · 74 cases this run.)
- **Grounding check** — Require every figure and factual claim in the answer to appear in — or be directly entailed by — the retrieved evidence; reject answers with unsupported spans (the model leaning on its own knowledge over the source). (Addresses: Meta leak, Misinterpreted, Out of scope, Table misread, Unit/format mismatch, Wrong period · 18 cases this run.)
- **Recompute check** — Recompute the arithmetic from the source figures and reject any answer whose stated result differs beyond a rounding tolerance. (Addresses: Math error, Misinterpreted, Out of scope, Unit/format mismatch, Wrong calc/quote · 18 cases this run.)
- **Numeric tolerance** — Normalize both values to a common precision and accept only within an explicit rounding tolerance (e.g. ± 0.5 of the least-significant digit). (Addresses: Misinterpreted, Out of scope, Table misread, Unit/format mismatch · 8 cases this run.)
- **Quote-vs-compute guard** — Decide whether the answer should be a verbatim quote or a computed value; validate a quote as a substring of the source and a computed value by recomputation. (Addresses: Wrong calc/quote · 2 cases this run.)

Flag — detect and route to review:

- **Cell-provenance check** — Confirm the answered value comes from the exact row + column the question names (row/column-header match); flag row/column drift. (Addresses: Out of scope, Table misread, Wrong entity · 5 cases this run.)
- **Direction/polarity check** — Compare the answer's polarity and direction (yes/no, increase/decrease, permitted/prohibited) against the source statement; flag reversals. (Addresses: deterministic:direction conflict · 4 cases this run.)
- **Evidence-coverage check** — Before answering, verify the retrieved context actually contains a span that supports the question; when coverage is below a set threshold, route to broader / multi-query retrieval or abstain rather than answering from a thin context. (Addresses: Table misread · 2 cases this run.)
- **Period-match check** — Extract the period/date from the question and require the answer's figure to come from that period; flag period mismatches. (Addresses: Out of scope · 1 case this run.)
- **Entity-match check** — Require the answer to reference the same entity named in the question (exact or alias match); flag cross-entity substitutions. (Addresses: Out of scope · 1 case this run.)
- **Abstention guard** — When evidence coverage is below a set threshold, require the model to abstain or hedge; flag confident answers made over weak or absent evidence. (Addresses: Certainty mismatch · 1 case this run.)

Case counts are per flagged answer and can overlap — a single answer may exhibit more than one failure pattern — so they need not sum to the number of failed questions. Patterns without a listed guard are semantic reasoning slips with no clean deterministic check; the remediation recommendations above address those.

9. Human oversight & sign-off

No human sign-off is on record for these exact results yet. The responsible person should review and sign the audit (recorded, hash-bound, in `audit_trail.jsonl`) before this card is published.

10. Reproducibility & provenance

Field	Value
Engine version	VERITROOPER v10
Run started	2026-08-26 15:27:11
Run finished	2026-08-26 17:29:18
Sampling temperature	0.1
Max tokens	1024
Results content hash (SHA-256)	f336a4ee6c66bbee...

System prompt given to the model under test:
Answer the question accurately using only the provided evidence.
Keep your response concise — one or two sentences.
Use only information from the evidence provided.

The run is reproducible from the stored record (certification_report.json + round_1.json + the generated question set) in the same package.

Limit of this evidence: replayability and repeated-run consistency show that the record and procedure can be checked; they do not by themselves establish broad validity. The questions are generated from the supplied material, scored denominators can differ when malformed items or connection failures are excluded, and configured model roles can introduce model-to-model variation. Generalization beyond the tested question set requires separate representative validation.

11. Data handling & isolation posture

VERITROOPER reads the evaluated material and the model without modifying or exporting either, and writes only its own audit artifacts, locally, into this results folder. When every selected endpoint (model under test and grader) is on your network or local, the audit runs with no outbound internet. See the accompanying **Data Handling & Isolation Posture** document for this run's recorded components and the honest scope of that claim.

12. Framework alignment

This audit produces evidence that **maps to** accuracy/robustness and documentation elements of the EU AI Act, the NIST AI RMF, and ISO/IEC 42001. It does not by itself establish conformity with any of them. The accompanying **Framework Crosswalk** shows the mapping control-by-control.

13. Deployment context & risk considerations

Provider-supplied — complete before publishing this card.

- Deployment context (where/how the system is used): _____
- Human-oversight measures in production (Art. 14): _____
- Known ethical / risk considerations & mitigations: _____

14. Card version

Field	Value
Card revision	this audit run (2026-08-26 17:29:18)
Supersedes	prior audits of the same panel (see the drift report, if any)

Generated by VERITROOPER as a system card documenting measured system behaviour. This artifact is testing evidence only. It does not by itself constitute, certify, or confer conformity with any framework, and does not replace the conformity assessment or the provider's other obligations. Have counsel review before external use.