

ALM TrustGuard

A trust-control layer that scores every AI-generated engineering artifact before a human is asked to trust it.

Students: Nguyen Huu Nguyen · Pham Huy Minh · **Academic supervisor:** Dr Ginel Dorleón · **Industry supervisor:** Khanh Nguyen (NEAX Co., Ltd.)
School: Science, Engineering & Technology, RMIT University Vietnam

01 The problem

An **Application Lifecycle Management (ALM)** platform is the system of record for engineering a regulated product — every requirement, test case and traceability link lives there, and audits are run against it. That platform now *writes* those artifacts with AI. In **ISO 26262** and **FDA-regulated** environments an unverified AI artifact is a liability — yet no engineer can review AI output at the speed AI produces it.

Hallucination

Invented conditions, omitted edge cases and ambiguous safety language flow straight into the ALM baseline — product defects, ASIL violations, recalls.

False compliance

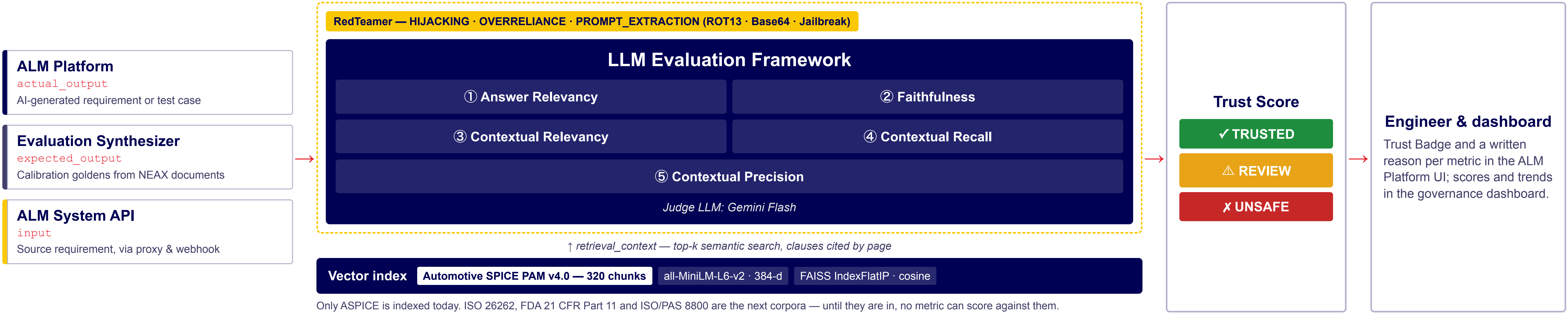
Artifacts pass superficial checks but are never measured against the ASPICE / FDA clauses they claim to satisfy — the gap surfaces months later, at review time.

Review bottleneck

Without a systematic trust signal, every artifact needs full manual review, which cancels out the entire productivity gain of using AI.

02 The solution — a non-invasive trust interceptor

TrustGuard sits between the ALM Platform and the engineer as a REST-API proxy and webhook listener. It **evaluates** AI output; it never generates or replaces it — and engineers change nothing about how they work.



03 How the score is built

#	Metric	What it detects	Weight
1	Faithfulness	Claims not grounded in retrieved standards — hallucinated conditions	0.30
2	Answer Relevancy	Output drifting off the source requirement	0.25
3	Contextual Relevancy	Retrieved compliance context is not applicable	0.20
4	Contextual Recall	Coverage gaps in the retrieved clauses	0.15
5	Contextual Precision	Most relevant clauses not ranked highest	0.10

```
Trust Score = 100 × ( 0.30 · Faithfulness
+ 0.25 · AnswerRelevancy
+ 0.20 · ContextualRelevancy
+ 0.15 · ContextualRecall
+ 0.10 · ContextualPrecision )
```

TRUSTED
80 – 100
Auto-approve and log

REVIEW
50 – 79
Flag + raise review ticket

UNSAFE
0 – 49
Block + request regeneration

05 Results — 200-artifact PHEV propulsion evaluation

200
AI-generated PHEV artifacts scored end-to-end

91.3
Mean Trust Score (median 93.7)

1.00
Mean Faithfulness — 197 of 200 fully grounded

176
Artifacts auto-approved as TRUSTED — zero blocked

The finding: switching from ASPICE process-level boilerplate to each requirement's own **traceability chain** as retrieval context lifted all three RAG metrics from near-zero (relevancy 0.02 → 0.94, recall 0.10 → 0.86, precision 0.05 → 0.85). Generation quality remained perfect: faithfulness 1.00, answer relevancy 0.95. **176 of 200 artifacts scored TRUSTED** and were auto-approved; 24 went to a human for review, and none were blocked. The right retrieval context, not a better generator, was the lever that unlocked compliance-grade trust scores.

06 Why it matters & what's next

For engineers
A Trust Badge appears inside the ALM Platform with a collapsible per-metric report. Review effort concentrates on the artifacts that actually need it — target: **≥ 30 % review-efficiency gain** at **< 3 s** per artifact.

For quality leads
Scores aggregate into a governance dashboard: which datasets, requirement types and standards clauses the AI is weakest on — a measurement, not an opinion, of where AI can be trusted.

Next
Expand the standards corpus beyond ASPICE (ISO 26262, FDA 21 CFR Part 11, ISO/PAS 8800), re-chunk for clause-level retrieval, calibrate thresholds on ≥ 200 goldens, and complete the red-team pass.