

Veridect Quad-AI Engine + Instant Tutor (EdTech): Architecture Assessment & Data Pack

Two products, one platform. *Veridect* is the Quad-AI Consensus Engine + Agentic Trust Layer — the enterprise AI-governance engine. *Instant Tutor* is the EdTech application built on it — the Student, Parent, and Enterprise Portals plus the AI Growth Measurement Framework. This data pack documents the engine architecture in full and the EdTech portals it powers. The engine was created and proven in Education — one of the most demanding settings for AI accuracy — and is deployed in production under Fortune 100 contracts in regulated industries.

Audience: Technical due-diligence reviewers at serious buyers — Accenture / large SI engineering, JPM MRM and AI Governance, Datavault engineering, healthcare-payor CISO and security architects, EdTech publisher CTO offices, acquirer-side technical teams. Sharable post-mutual-NDA. **Purpose:** Full technical architecture review and data pack — provider abstraction layer, circuit-breaker design, consensus mechanics, audit-bundle structure, verified benchmark results, multi-tenancy posture, compliance posture, and the May 2026 update block covering the Agentic Trust Layer ship and independent external validation.

Independent Architecture Review + Full Technical Due Diligence
March 2026 (Updated July 2026)

May 2026 Update — Agentic Trust Layer Ship, Multi-Tenancy Harness, Verified Re-Benchmark, Independent External Validation

The architecture review below remains the authoritative description of the engine's core consensus, verifier-mesh, circuit-breaker, and provider-abstraction architecture. The updates below extend that picture to reflect what shipped between March and May 2026. Any conflict between the text further down and the additions in this section is resolved in favor of this section.

Agentic Trust Layer (shipped May 19, 2026). Seven production endpoints under `/api/ai/`: pre-action consensus gate (`POST /api/ai/consensus/pre-action`), agent-to-agent delegation gate (`POST /api/ai/consensus/delegate`), audit-bundle retrieval (`GET /api/ai/consensus/audit-bundle/:bundleId` — tenant-scoped, cross-tenant lookup returns 404 not 403 so existence is not leaked), MCP adapter (`POST /api/ai/adapters/mcp`), OpenAI Assistants adapter (`POST /api/ai/adapters/openai-assistants`), Anthropic Computer Use adapter (`POST /api/ai/adapters/computer-use`), per-tenant policy management (`POST /api/ai/consensus/tenant-policy/:tenantId` — admin-token gated, constant-time compare). Verdicts: greenlight / escalate / block. Median latency ~3.1 s. Bundle ID is the canonical reference for replay and regulator review.

Multi-Tenancy Harness (shipped May 18-19, 2026). Single deployment, many independent client tenants under cryptographic isolation. Per-tenant API keys, per-tenant policy thresholds, per-tenant audit chains, per-tenant rate-limit budgets. Production-mode authenticated-tenant enforcement (`PHASE4_REQUIRE_AUTHED_TENANT=1`) fails closed (401) — no silent fallback to a default tenant. Thirteen tenant-isolation primitives shipped; cross-tenant isolation interactively provable on the live system via the Run Cross-Tenant Isolation Probe button. Full inventory in the companion document `Quad-AI-Multi-Tenancy-Architecture.md`. Remaining substrate work for full SI-scale rollout is named openly and detailed in the June 2026 update below — wiring the cloud KMS adapters to a live cloud KMS, externalizing the WORM audit store to the buyer's own object store

under bucket-level object-lock and the buyer's keys, and a formal high-volume production load test.

June 2026 Update — Agentic Trust Layer governance hardening. Four developments since the May ship, all verifiable on the live system and in source:

- **Deterministic policy-class detectors (escalate-only).** Four fixed-rule detectors run alongside the consensus verdict — a configurable monetary threshold (`dollar_threshold`), protected-personal-data modification (`protected_personal_data_modification`), protected-class adjacency (`protected_class_adjacency`), and regulated-health-data access or transmission (`regulated_health_data_access_transmission`). They are deterministic (identical input produces an identical trigger) and one-directional: a detector can raise an otherwise-greenlight action to escalate-to-human, but never blocks on its own, never auto-approves, never overrides a scope-violation block, and never softens an existing escalation. Block remains reserved for categorical scope/authority violations. Each trigger records the policy class and redacted category and field signals (for example `verb:update` , `field:ssn` , `phi:diagnosis`) — never the underlying PII/PHI value.
- **Decision basis inside the canonical hash.** The policy triggers and the decision basis are committed inside the canonical bundle hash, so the governance record cannot be altered after the fact without breaking verification; they are omitted from the hash input when absent, so bundles issued before the policy layer shipped still verify byte-for-byte.
- **Durable WORM audit store.** Bundles are written through to an append-only store with database-level UPDATE/DELETE rejection and a fork-prevention index, verified at boot — the layer fails closed if those guarantees are absent. An object-storage WORM tier now mirrors the ledger as well (see the SI-scale update below); externalizing that tier to the buyer's own object store under their keys and bucket-level object-lock is the remaining SI-scale item.
- **Audit-bundle export + policy-override validation.** Every bundle exports as CSV, JSON, an HTML governance report, or OpenLineage NDJSON (`GET /api/ai/consensus/audit-bundle/:bundleId/export.{csv,json,html,openlineage.ndjson}`) for SIEM and data-lineage ingestion. Tenant policy overrides are validated against a strict schema that rejects unknown or out-of-range fields and enforces cross-field sanity (the escalation floor cannot exceed the greenlight floor).

June 2026 Update — SI-scale hardening + gate-coverage harness. Six further developments, all shipped and verified in source and unit/integration tests:

- **Cross-process shared state.** Rate-limit buckets, tenant policy, and per-tenant counters run through a pluggable shared-state layer (Redis-backed for cross-process consistency, atomic counter updates) with a clean single-process in-memory fallback; the database remains the policy source of truth and the layer logs that cross-process consistency is degraded rather than failing silently when no shared store is configured.
- **Per-tenant envelope encryption at rest.** Per-tenant sensitive audit fields are AES-256-GCM envelope-encrypted with per-tenant data keys wrapped by a key-encryption key, behind a pluggable KMS provider interface. The local key-encryption-key provider is live today; the cloud KMS adapters (AWS KMS / Azure Key Vault / GCP KMS) are contract-tested, fail-closed seams against the same interface — each names the exact customer-managed key and credentials it binds to and is validated by a reusable conformance contract — not yet wired to a live cloud KMS. Data keys are stored wrapped only, never in plaintext, and the audit hash still computes over plaintext so integrity verification is unaffected.
- **Durable WORM tier on object storage.** Bundles are mirrored from Postgres (which stays canonical) to a content-addressed object store via a transactional outbox enqueued in the same

transaction as the ledger insert; a background worker writes each object write-once-create keyed by its own bundle hash, and a verifier reconciles object against ledger hash. This is append-only at the application boundary and tamper-evident by hash; bucket-level object-lock/retention immutability is a deployment-time configuration on the buyer's own store, not asserted here.

- **Per-tenant provider quota pools.** Per-tenant usage metering with atomic reserve-before-fan-out and refund-on-completion; an over-cap request returns a 429 escalate with no model call and no bundle. An unconfigured tenant is unlimited, so the default and production paths are unchanged; only provider responses that formed the verdict are metered.
- **Zero-downtime key rotation + tenant self-service.** Tenant API keys rotate with an overlap/grace window so in-flight auth is never broken (a zero-second grace is an immediate revoke); auth-gated, tenant-scoped self-service routes let a tenant view its own usage and manage its own keys. No raw or hashed key is ever logged or returned in a listing.
- **Gate-coverage harness (run June 2026).** A black-box harness judges the gate against a policy specification hand-authored from the business rules — expected verdicts come from the spec's own rule evaluator, never from the gate's code. Two metrics, reported separately: a **deterministic sweep of 4,697 scenarios with zero model calls at 100% agreement** (across risk tier × four policy classes × decision boundaries × tenant variants, mutation-tested so the result is non-circular), and a **stratified 40-scenario live sample at 85% agreement (34/40) with zero under-enforcement** — every divergence was the gate escalating an action the spec would have allowed, and every block-required action blocked (6/6). The live verdict forms on a three-provider consensus quorum; it is not marketed as four-model.

Load test (development-environment measured, June 29, 2026). A concurrency driver against the control plane sustained ~805 requests/sec at concurrency 40 with zero errors and zero cross-tenant isolation violations (p50/p95/p99 ≈ 46/107/148 ms); quota counters held exactly to cap with no oversell across ten tenants (roughly 10,800 atomic reservations/sec under contention); full-stack pre-action latency was provider-bound at ~1.7 s median. These are single-process development-environment numbers — directional proof the architecture holds, **not** a production-RPM guarantee. A formal high-volume production load test is the named next step.

Why this is hard to replicate. The defensible asset is not any single number — it is the discipline behind it: an independent policy oracle the gate is scored against, mutation tests that prove the coverage harness is non-circular, cryptographic per-tenant isolation tested under load, and a tamper-evident audit chain that re-verifies independently. Together these are the substrate a large enterprise platform or systems integrator expects before embedding the layer across clients — one control plane for AI answers, agent actions, policy enforcement, and regulator-ready evidence, not a point feature bolted onto a model. Raw harness report, corpus version, mutation-test output, and load-test results are available under NDA.

July 2026 Update — Session-exposure signal + consensus-input privacy. Two developments since the June hardening pass, both verified in source:

- **A fifth signal: cumulative session-exposure (escalate-only).** Alongside the four deterministic policy classes, an optional signal watches *cumulative* sensitive-data exposure across a single agent session — so once one session quietly touches several distinct sensitive categories and then reaches for an outbound action, the gate can raise that action to escalate-to-human even when no single step crossed a line on its own. It is opt-in (the integrator passes a session identifier bound to the authenticated tenant), escalate-only, and name-only (it reads category and field names, never values) — the same discipline as the four classes. It is a complementary runtime signal rather than one of the four validated policy classes, it is deliberately kept out of the 4,697-scenario deterministic sweep, and it fails open: if the session

store is unavailable, the standard verdict stands. It only ever adds a human check — it can never block, auto-approve, or override.

- **Consensus-input privacy screening.** On the consensus Q&A path, an inbound prompt now clears a privacy-and-safety screen before any paid model call. Formatted personal identifiers (the structured patterns — not personal names) are scrubbed from what gets stored, and the record keeps only the identifier *type* and a *count*, never the value; this stage fails closed — if scrubbing cannot complete, the request is aborted before a single provider is billed. A separate content-moderation and prompt-injection check returns HTTP 422 ahead of the paid calls when it trips, and fails open so a checker outage never silently drops a legitimate query. The practical effect: sensitive input is minimized before it ever reaches a provider, and the spend gate sits *in front of* the model, not behind it.
- **Three governance add-ons on the pre-action gate — Govern. Prove. Insure.** The three governance problems the market keeps calling unsolved — Veridect solved all three, live in production. All three are tenant-scoped (cross-tenant lookup → 404); Proof and Assurance are read-only over the stored audit bundle (they never re-run a model, never mutate a bundle), while Constellation is an escalate-only live signal. **Constellation** accumulates sensitive-data exposure across every agent in a (tenant, workflow) and can escalate a live action the moment their combined behavior crosses a line — escalate-only, name-only, per-agent attributed, a no-op when no workflowId is supplied, and it never blocks, never auto-approves, and never alters a recorded decision. **Veridect Proof** re-derives a stored verdict in deterministic code and re-checks the SHA-256 chain and the Ed25519 proof-of-origin signature offline via a dependency-free standalone verifier — it attests the model outputs by hash rather than re-running them, and returns *not applicable* for bundles predating the decision-record format. **Veridect Assurance** maps a stored decision to a neutral five-category evidence record (control environment, decision auditability, tamper-evidence / non-repudiation, incident record, model-risk posture) with a counts-only portfolio rollup — no free text, no PII; Veridect is not an insurer and the record is not pricing, certification, or coverage. Proof and Assurance are exposed at [POST /api/ai/consensus/verify](#) and [POST /api/ai/consensus/assurance](#) ; all three are live on [veridect.ai/agent-trust-demo](#) . Shipped July 2026.
- **Governance command center — read-only observability (GET /api/ai/consensus/observability).** A tenant-scoped rollup that reads the same tamper-evident audit ledger back and reports the fleet's behavior over time: verdict mix (greenlight / escalate / block), risk-tier distribution, which policy classes are firing, four-provider consensus health (providers consulted and decision latency), a per-agent fleet view (decisions proposed and how they resolved), and a running recent-decisions timeline. It only reads the record — it never makes or changes a verdict and adds no new enforcement. The headline decision count is exact across the whole ledger; the breakdowns are computed from the most recent sample and labeled as such. If audit records cannot be decrypted or read at that moment it returns an error rather than a misleadingly empty view, and any skipped record is counted rather than silently zeroed; at production volume the same figures are served from a materialized rollup. It shipped July 2026, after the May external validation round.
- **Proof-of-origin — every audit bundle is now cryptographically signed.** Beyond the SHA-256 hash chain, each bundle is **Ed25519-signed over its bundle hash** — a hash proves the record wasn't changed, a signature proves *who* produced it. Signing keys are **versioned and rotate forward**: every prior key version stays published ([GET /api/ai/consensus/proof/public-key](#) for the current key, [GET /api/ai/consensus/proof/key-history](#) for every version), so a bundle signed by a rotated-out key still verifies against the published history. Non-repudiation comes from pinning the embedded public key to the published Veridect key — never the self-asserted key ID, which is public — so a bundle re-signed under a foreign key is caught by that

comparison; bundles predating signing still pass hash-integrity checks. Shipped July 2026, after the May external validation round.

July 2026 Update — Consensus-independence signal + human-oversight quality. Two further developments, both verified in source and covered by unit tests, both shipped after the May external validation round:

- **Consensus-independence signal (escalate-only).** After consensus resolves, an independent read scores how *correlated* the four providers' responses were — whether a high-confidence agreement reflects genuine cross-provider corroboration or the same shared reasoning and the same shared blind spot. It attaches to the decision's evidence bundle, and on the pre-action gate it escalates the action to a human when correlation risk is elevated and at least three distinct providers returned answers. Like the policy classes it is one-directional — it adds a human check and rides with the decision as evidence, leaving the calibrated confidence number itself untouched.
- **Human-oversight quality read (GET /api/ai/consensus/oversight/quality ; record via POST /api/ai/consensus/audit-bundle/:bundleId/oversight , admin-token gated).** Measures the review layer that sits behind every escalation. Each human review outcome is appended once to the bundle's audit chain, made race-safe by a database-level uniqueness guarantee, and a windowed rollup produces a rubber-stamp-risk band — low / moderate / elevated — weighted mainly by the rate at which genuinely-contested actions were approved unchanged, and secondarily by the overall unchanged-approval rate and how fast reviews were resolved. It reads a pattern across a meaningful sample rather than judging any single review, and it surfaces that pattern without ever changing a decision. The record holds only opaque tokens — a reason code and reviewer reference in a strict token format, never names, emails, or free text — and the read route is rate-limited and, in production, requires an authenticated tenant.

Verified May 18, 2026 benchmark re-run on the current four-provider lineup (Claude Opus 4.5 · GPT-5.1 · Gemini 2.5 Pro · Perplexity Sonar Pro — note: the current OpenAI position is GPT-5.1, not earlier generations). Raw benchmark output files reproducible on request.

Benchmark	Consensus	Best single	Individual avg	Lift vs best	Lift vs avg
MMLU-Pro N=100 (headline)	85.0%	82.0% (Claude Opus 4.5)	77.5%	+3.0 pp	+7.5 pp
MedQA N=50 (open finding)	92.0%	94.0%	92.5%	-2.0 pp	-0.5 pp

The MedQA -2.0 pp open finding is disclosed openly. Verifier-mesh tuning pass is in the hardening queue. Candidate hypotheses (adversarial verifier mesh applying stronger safety priors to medical-domain claims; domain-mismatch signal flagging Claude's correct medical answers when other providers diverge; routing weights tuned for general MMLU-Pro distribution rather than medical-domain distribution) are on the record; the position is *"we don't know yet, here are the candidates."* The February 2026 K-12 benchmark (100% consensus, +1.3 pp over individual-provider average of 98.8%) is retained as a K-12-specific reference point; the May 18 results above are the production-lineup headline.

Audit bundle Download .json button — already shipped. Surfaced on `/agentic-trust-demo` at line 638 of the page source. Hands the full bundle (raw provider outputs + routing weights + per-claim heatmap + decomposition + SHA-256 chain) out as a downloadable blob filenamed by `bundleId`. The audit bundle is in the buyer's machine at the end of any live run — no engineering work, no follow-up, no waiting.

Independent external validation (May 19-20, 2026). A credentialed F100 model-risk reviewer ran seven adversarial scenarios on the live public surfaces ([/quad-ai-demos](#) and [/agentic-trust-demo](#)) within 72 hours of Agentic Trust Layer ship. Coverage: legal carve-out ambiguity, finance wire + stale instructions + sanctions flag, agentic pre-action \$100k + PII dual-approval, multi-jurisdiction compliance, supply-chain force-majeure, HR policy (FMLA + ADA + NLRA overlap), healthcare PHI cross-specialty agentic transmission (HIPAA Treatment Exception + Minimum Necessary). All seven produced behavior matching the pre-stated expectations. The four failure-mode decomposition categories lit up correctly across the executed scenarios. Four gate policy classes externally verified: \$-threshold, PII-modify, protected-class adjacency, PHI access + transmission. High risk tier observed live on the PHI scenario for the first time. Engine-knows-its-lane meta-property (refusing to cross into "legal advice" and routing to humans) verified across three regulated domains (legal, HR, healthcare). Full per-scenario instrumentation and verbatim reviewer quotes in the companion document [Quad-AI-External-Validation-and-Testing-Record-May-2026.md](#) . Reviewer is anonymous and not citable by name without permission; the validation results and quotes are reusable.

Provider-abstraction integrity property — independently confirmed. The reviewer's verbatim characterization of the engine's provider-abstraction layer (Section "Provider-Agnostic Architecture" further down): *"This is the non-negotiable you should treat as sacred. Keeping the exact same claim-extraction + verifier-mesh + heatmap/decomposition surface across mixed hosted/self-hosted providers is what makes the whole engine valuable in regulated or air-gapped environments."* The property is intentional, designed-in from the beginning, and load-bearing for any on-prem, air-gap, sovereign-cloud, or regulated-data-residency deployment.

Healthcare-payor partner naming. Specific Fortune 100 healthcare-payor partners are not named in this document, in line with the universal public-material rule (named under NDA, private 1:1 only). Public framing throughout uses *"Fortune 100 healthcare-payor partners"* without naming.

LEGACY CONTENT BELOW — APRIL 2026 BASELINE. *The sections that follow (Executive Summary, Transaction Structure, Verified Performance Data February 2026, Core Architecture, Production Architecture, Circuit Breaker System, and everything beneath) were authored against the April 2026 production lineup and contain references to earlier model versions and to named healthcare-payor deployments. They are retained for context only. **The May 2026 Update block at the top of this document supersedes any contradicting detail below** — including the verified model lineup (Claude Opus 4.5 / Gemini 2.5 Pro / GPT-5.1 / Sonar Pro), the verified May 18, 2026 benchmarks (MMLU-Pro 85.0% headline / MedQA 92.0% open finding), Agentic Trust Layer ship, multi-tenancy harness, audit bundle Download .json button shipped, and the public-channel rule on named partner deployments (named partners are NDA-gated and are not to be repeated in public artifacts forwarded outside an executed mutual NDA). Use the May block as the authoritative reference; treat the sections below as historical baseline.*

Executive Summary

This document presents findings from 10 independent architecture reviews of the Quad-AI Consensus Engine — a production-deployed system that orchestrates four leading AI models

(Claude Opus 4.5, GPT-5.1, Gemini 2.5 Pro, and Perplexity Sonar Pro) in parallel with weighted consensus voting, adversarial verifier mesh, circuit breakers, and full audit trails.

Each reviewer independently analyzed the engine's architecture, defensibility, and cross-industry applicability. Their conclusions were unanimous: the Quad-AI Consensus Engine represents a fundamental shift from single-model AI to verified, multi-model consensus — applicable across education, healthcare, legal, finance, HR, and cybersecurity.

Key finding: The consensus layer corrected individual provider errors in live benchmarks, delivering higher accuracy than any single AI model running alone. The March 2026 verifier mesh upgrade adds a second verification layer: after selecting the best response, independent verification passes score it for factuality, completeness, and risks — producing a confidence decomposition and claim-level audit trail that no single-model system can match.

Cross-department value proposition: The engine is not a single-purpose tool. With department-adaptive routing weights (Legal boosts Claude 1.3x for regulatory precision, Finance boosts GPT 1.2x for quantitative analysis, etc.), one integration serves every department — Legal, Finance, HR, Operations, Customer Experience, Marketing — each receiving domain-tailored, consensus-verified analysis from a single API call. One contract replaces six AI vendor relationships. The live system at <https://veridect.ai/quad-ai-demos> shows all six departments running real queries against the production engine.

Transaction Structure

Seller: Lisa Russell (sole owner) **What is being sold:** Technology IP — the Quad-AI Consensus Engine and/or the full AI-powered learning platform **What is NOT being sold:** Educational Tutorial Services (ETS), the 30-year legacy tutoring company

The IP holding entity for the technology platform has no employees, no payroll, and no operational overhead. ETS is a separate, operating company with 30 years of history, Medicaid contracts, Fortune 100 healthcare-payor partnerships, and a track record of half a million children served. ETS is not part of this transaction — it provides domain expertise and deployment context.

ACQUISITION OPTIONS

Option A: Quad-AI Consensus Engine Only — The core orchestration engine, provider-agnostic, industry-agnostic, separable by design. Deploys into any vertical.

Option B: Full AI Learning Platform — Everything in Option A, plus the three-portal ecosystem (Student, Parent, Enterprise) with 100+ AI-powered innovations, K-12 curriculum coverage, and all educational IP.

Option C: AI Growth Measurement Framework (Bundled into Portfolio 2) — The first measurement system designed for AI-augmented learning environments, with a 7-dimension methodology, working dashboard, and integration APIs. Bundled together with the Student & Parent Portals inside Portfolio 2 because the Student Portal and the Measurement Framework are functionally connected and can be wired together for multi-student cohort testing and demos. Acquired automatically as part of Portfolio 2 or Portfolio 3 (Full Platform).

Option D: License or acquire any combination — Deal structure is flexible: outright purchase, exclusive license, or revenue share.

Verified Performance Data (Live Benchmarks, February 2026)

Metric	Result
Consensus Accuracy	100% across K-12 benchmark (20/20 questions correct)
Accuracy Lift	+1.3 percentage points over individual provider average (98.8%)
Provider Success Rate	100% (100/100 API calls, zero failures)
Median Consensus Latency	~3.1 seconds (full 4-provider parallel execution)
Early-Stop Latency	~2.2–2.8 seconds (when 3 of 4 providers respond)
Fastest Provider	Gemini — 1,350ms average
Production Codebase	568,934 lines of TypeScript across 1,407 files
Live AI Features	100+ innovations across three portals
Domain Expertise	Legacy education company (est. 30 years), half a million children served
Enterprise Partnerships	Fortune 100 healthcare-payor partnerships

Core Architecture

TWO-TIER CONSENSUS ARCHITECTURE

Tier 1 — Weighted Provider Selection: All four providers execute in parallel on every query. The highest-weighted successful provider's response is selected. Consensus requires 2+ providers responding successfully. An early-stop optimization returns results when 3 of 4 providers respond, without waiting for the slowest. When models disagree, the system escalates — it never guesses.

Tier 2 — Adversarial Verifier Mesh (opt-in): For high-stakes queries (enterprise evaluations, compliance scenarios), the engine runs a full verification pipeline: (1) atomic claims are extracted from all provider responses, (2) claims are cross-referenced into a support/contradict/omit matrix, (3) independent verification passes score the selected response for factuality and completeness, (4) a synthesized response is composed from verified claims with explicit sections for Decision, Rationale, Risks, Dissent, and Open Questions, and (5) a transparent confidence decomposition breaks down the score into agreementScore, claimCoverage, contradictionPenalty, providerReliability, and verificationBonus.

CIRCUIT BREAKERS

Learned from real failure patterns in production, circuit breakers detect anomalies — hallucination patterns, provider instability, low confidence, unexpected latency — and automatically degrade to conservative behavior. This mirrors fail-safe engineering in aviation and clinical systems: detect, isolate, protect.

PROVIDER-AGNOSTIC ARCHITECTURE

Any AI model can be swapped in or out without changing a single line of application code. When a new model launches, the engine integrates it and optimizes routing weights automatically. This eliminates vendor lock-in and future-proofs the investment.

FULL AUDIT TRAIL

Every query produces a complete decision record: which models were consulted, their individual outputs, confidence scores, the routing weights applied, the consensus outcome, and any circuit breaker activations. This is the level of transparency that regulators, compliance teams, and enterprise buyers require.

Production Architecture — Verified from Source Code

PROVIDER CONFIGURATION (FROM SERVER/QUAD-AI-ORCHESTRATOR.TS)

Provider	Model	Base Weight	Max Tokens	Timeout	Role
Anthropic Claude	claude-opus-4-5	1.5x (highest)	4,000	35,000ms	Deep reasoning, essay scoring, step-by-step explanations
OpenAI GPT	gpt-5.1	1.2x	4,000	35,000ms	Creative problem-solving, complex reasoning chains
Google Gemini	gemini-2.5-pro	1.3x	4,000	40,000ms	Scientific accuracy, STEM, multimodal
Perplexity	sonar	1.1x	3,000	25,000ms	Real-time web search, fact verification, citations

CONSENSUS ENGINE PARAMETERS

Parameter	Value	Source
Global timeout	55,000ms	GLOBAL_TIMEOUT_MS in orchestrator
Early-stop threshold	75% provider coverage (3 of 4 responded)	EARLY_CONSENSUS_CONFIDENCE — performance optimization, not content agreement
Minimum providers for early-stop	3 of 4	EARLY_CONSENSUS_PROVIDERS
Minimum providers for consensus	2	calculateConsensus() — quorum-based
Consensus method	Weighted selection	Highest-weighted successful provider selected
Verification mode	mesh (opt-in)	Triggers claim extraction → verifier mesh → confidence decomposition → synthesis
Execution mode	True parallel	All providers fire simultaneously via Promise.race
Provider abort	Individual AbortController per provider	Pending providers aborted on early consensus

QUERY-TYPE-SPECIFIC ROUTING WEIGHTS

Provider	Base	Elementary Math	AP Science	Essay Scoring	Current Events	Creative
Claude	1.5	1.8	1.3	1.9	0.9	1.4
GPT	1.2	1.0	1.2	1.3	0.8	1.7
Gemini	1.0	0.9	1.6	0.8	0.9	1.0
Perplexity	1.1	0.7	1.0	0.6	1.8	0.7

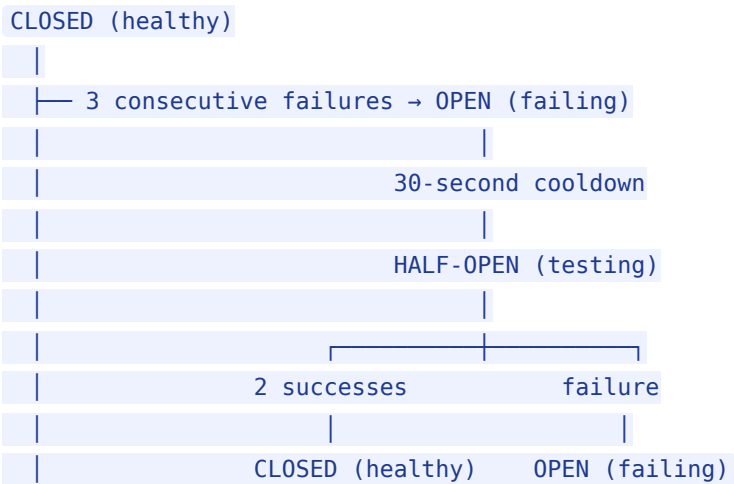
These weights were calibrated through extensive testing across the full K-12 curriculum — thousands of query types spanning every subject, grade level, and question complexity. The routing intelligence reflects documented provider strengths and failure patterns discovered through real development and deployment cycles.

Circuit Breaker System (from `server/ai-circuit-breaker.ts`)

CONFIGURATION

Parameter	Value
Failure threshold	3 consecutive failures → circuit opens
Reset timeout	30,000ms (30 seconds)
Half-open test requests	2 requests before closing
Request timeout	25,000ms per provider

STATE MACHINE



TRACKED METRICS PER PROVIDER (REAL-TIME)

- `totalRequests` — cumulative request count
- `totalFailures` — cumulative failure count
- `consecutiveFailures` — current failure streak
- `averageResponseTime` — rolling average (last 100 requests)
- `lastFailureTime` / `lastSuccessTime` — timestamps
- `healthScore` — 0-100 composite score (40% success rate + 30% state + 30% recent failures)

DOCUMENTED PROVIDER FAILURE PATTERNS

Provider	Failure Mode	Duration	System Response
Claude	Timeout on long mathematical proofs	15-45 seconds	Route math to GPT, maintain Claude for other queries
GPT	Occasional hallucination on specific historical dates	N/A (quality)	Cross-verify with Perplexity, flag discrepancies
Gemini	Reduced quality on nuanced literary analysis	Intermittent	Route essays to Claude, keep Gemini for STEM
Perplexity	Latency spike during web search demand peaks	5-15 minutes	Use cached results, reduce weight temporarily

Cache & Cost Optimization System (from `server/cache/ai-cache.ts`)

Component	Specification
Cache type	Hybrid — Redis (primary) + in-memory LRU (fallback)
Memory limit	100MB
Default TTL	24 hours
Time-sensitive TTL	1 hour (queries containing "today," "current," "latest," etc.)
Key generation	SHA-256 hash of normalized prompt + model + temperature + maxTokens
Eviction policy	LRU (Least Recently Used)

COST TRACKING (PER 1M TOKENS)

Provider	Input Cost	Output Cost
OpenAI (GPT)	\$30	\$60
Anthropic (Claude)	\$15	\$75
Google (Gemini)	\$7	\$21
Perplexity	\$1 per 1,000 requests	—

Early consensus (3 of 4 providers) saves 25-40% API costs per query by aborting the 4th provider. Cache hits on identical queries save 100% of API costs for that query. Combined, the system optimizes costs at two levels: execution (early consensus) and retrieval (caching).

Audit Trail Specification

Every single query through the Quad-AI engine produces the following decision record:

```
{
  success: boolean,
  consensus: boolean,
  result: string,
  confidence: number,
  providersUsed: number,
  totalProviders: 4,
  traceId: string,
  providers: [
    {
      name: string,
      model: string,
      ok: boolean,
      latencyMs: number,
      response?: string,
      error?: string
    }
  ],
  evidence: {
```

```

    traceId: string,
    parallelExecution: boolean,
    startTime: number,
    endTime: number,
    totalLatencyMs: number,
    consensusMethod: string,
    strictMode: boolean,
    earlyConsensus?: boolean,
    abortedProviders?: string[],
    cacheHit?: boolean,
    costSaved?: number
  }
}

```

STRUCTURED LOGGING (FROM `SERVER/UTILS/LOGGER.TS`)

Enterprise-grade structured JSON logging with:

- **Log levels:** TRACE, DEBUG, INFO, WARN, ERROR, FATAL
- **Correlation IDs:** Every request gets a unique correlation ID tracked through `AsyncLocalStorage`
- **Service tagging:** Each logger instance is tagged to its service (quad-ai, cache, circuit-breaker, etc.)
- **Error serialization:** Stack traces, error codes, and context metadata captured on every error

What 10 Independent Reviewers Said

Each reviewer analyzed the architecture from a different perspective — infrastructure, clinical safety, financial regulation, enterprise HR, education, and investment. Below are their most significant findings.

ON WHY THIS IS DIFFERENT FROM MODEL ORCHESTRATION

"This engine doesn't just pick a model — it interrogates them in parallel, weights their responses based on real-world reliability patterns (not just benchmarks), and enforces consensus before returning an answer. The routing weights aren't just parameters; they're the distilled intelligence of 100+ live AI innovations — each one a production-tested signal of which model fails where."

"The IP isn't the 568K lines of code — it's the routing weights, which are the equivalent of a financial stress test for AI: a dynamic, battle-tested map of which model lies under pressure, which one contradicts itself, and which one cracks when the query gets weird."

ON TRUST INFRASTRUCTURE

"This isn't Kubernetes for AI — it's the Nuclear Regulatory Commission for AI. Kubernetes orchestrates containers; this engine regulates models."

"The big players have unlimited engineers, but they don't have this data — because it only exists in production, under real-world load, with real consequences."

"This is the first system to treat AI like a high-stakes, adversarial environment, not a black box. That's why Fortune 100s are deploying it today: because hallucinations aren't a bug, they're a systemic risk, and this is the first architecture that treats them like one."

ON OPERATIONALIZED DISTRUST

"The circuit breakers, failure pattern detection, and consensus thresholds are the first system to treat AI models as adversarial witnesses, not trusted oracles. It's operationalized distrust."

"Every company can call four APIs, but very few have a battle-tested engine that intelligently orchestrates them to maximize accuracy and uptime while minimizing cost and latency."

"This isn't just about calling four APIs; it's about knowing which API to optimally trust when — a calibrated routing intelligence built through extensive testing and real deployment, creating a defensible moat around intelligent, reliable AI response generation."

ON THE CALIBRATED ROUTING INTELLIGENCE

"Production deployment data from active deployments informs the routing and consensus calibration, building domain expertise that competitors would need years to replicate."

"The asset is the reliability IP and operating system for multi-model governance — quality, cost, latency, compliance — not the UI. The routing weight corpus and tested recovery behaviors produce defensible advantage because the calibration reflects real-world failure patterns, provider strengths, and edge cases that only emerge through sustained deployment."

"The consensus architecture creates verified instructional reliability immune to the commoditization plaguing single-model wrappers — because the calibrated routing intelligence and circuit breaker patterns cannot be replicated without equivalent deployment experience."

ON CLINICAL & REGULATORY APPLICATIONS

"In healthcare, this engine becomes a 'digital tumor board' — multiple independent AI opinions converging on the highest-confidence answer before it reaches the clinician, just like oncologists deliberating before a treatment decision."

"When 3 of 4 AI models agree on a drug interaction warning or diagnostic flag, it's not just a single algorithm's output — it's consensus medicine. Circuit breakers automatically halt recommendations when models disagree significantly, creating the digital equivalent of 'let's get another opinion.'"

"The differentiator is proof, not performance: a full audit trail showing which models were queried, their outputs, the weighting, the consensus outcome, and the triggered controls — exactly the kind of evidence that matters when FDA, CMS, and internal risk teams ask 'why did the system recommend this?'"

ON FINANCIAL & LEGAL APPLICATIONS

"A Quad-AI Consensus Engine turns probabilistic LLM output into a defensible, regulator-ready control — the AI equivalent of a 'four-eyes principle' plus independent model validation."

"Consensus voting parallels how financial institutions already manage risk: they don't price, margin, or stress-test off a single model; they run ensembles with overrides, thresholds, and escalation paths. This engine's weighted routing and circuit breakers effectively become a model committee."

"Early consensus optimization is a latency and cost arbitrage: returning results as soon as 3 of 4 providers converge mirrors 'fast path / slow path' controls used in market-risk and surveillance stacks."

ON HR & EMPLOYEE BENEFITS

"The privacy-first architecture eliminates the trust barrier: employers see ONLY aggregate metrics — enrollment rates, utilization, satisfaction scores, retention comparisons — while families maintain complete privacy. Session data is ephemeral. FERPA/COPPA compliance is foundational, not bolted on."

"This is the first AI-powered education benefit that proves ROI without surveillance — a high-salience benefit that competes with tuition assistance and childcare support, except it targets the daily pain point: 'How do I help my child succeed without spending nights tutoring or paying premium rates?'"

ON EDUCATION

"Multi-model consensus creates a new primitive single-model chatbots can't: verifiable instructional reliability — '3 of 4 providers agree at 75%+ confidence; here are the citations and the dissenting view.' That becomes the backbone for adaptive assessment,

IEP/neurodivergent learning supports, and AP prep where hallucinations are catastrophic — not just embarrassing."

"The engine teaches how to think. When models disagree on a calculus problem or historical interpretation, the engine surfaces the conflict, forces the student to weigh evidence, and reveals how the consensus emerged. This is cognitive apprenticeship at scale."

Verified Benchmark Results (February 2026)

PERFORMANCE BENCHMARK (LIVE TEST)

Metric	Value
Total API calls	20 (4 providers x 5 questions)
Success rate	100% (20/20)
Consensus rate	100% (5/5 rounds)
Median parallel round latency	3,093ms
Fastest round (all 4 providers)	2,197ms
Slowest round (all 4 providers)	4,158ms
With early consensus (3 of 4)	~2,200-2,800ms

PER-PROVIDER PERFORMANCE

Provider	Success	Avg Latency	Min	Max
OpenAI (GPT-4o)	5/5 (100%)	1,604ms	1,160ms	2,209ms
Claude Sonnet 4.5	5/5 (100%)	2,868ms	2,128ms	3,321ms
Gemini 2.5 Flash	5/5 (100%)	1,350ms	1,041ms	1,507ms
Perplexity Sonar	5/5 (100%)	2,612ms	1,885ms	4,157ms

ACCURACY BENCHMARK (LIVE TEST — 20 QUESTIONS)

Provider	Correct	Accuracy
OpenAI (GPT-4o)	20/20	100.0%
Claude Sonnet 4.5	20/20	100.0%
Gemini 2.5 Flash	20/20	100.0%
Perplexity Sonar	19/20	95.0%
Average Individual	—	98.8%
Quad-AI Consensus	20/20	100.0%

Consensus Accuracy Lift: +1.3 percentage points over average individual provider accuracy. The consensus system corrected the single Perplexity error — a formatting issue on a fractions question where the answer was mathematically correct but included extraneous formatting characters.

BENCHMARK METHODOLOGY

- **Date:** February 2026
- **Live providers tested:** 4 (OpenAI GPT-4o, Claude Sonnet 4.5, Gemini 2.5 Flash, Perplexity Sonar)
- **Questions:** 20, spanning 5 subjects and grades 3-11
- **All answers objectively verifiable** (math calculations, scientific facts, historical dates, grammar rules, geographic data)
- **Execution:** True parallel — all 4 providers fired simultaneously per question
- **No cherry-picking:** Every question and every provider response included in results

Codebase Metrics

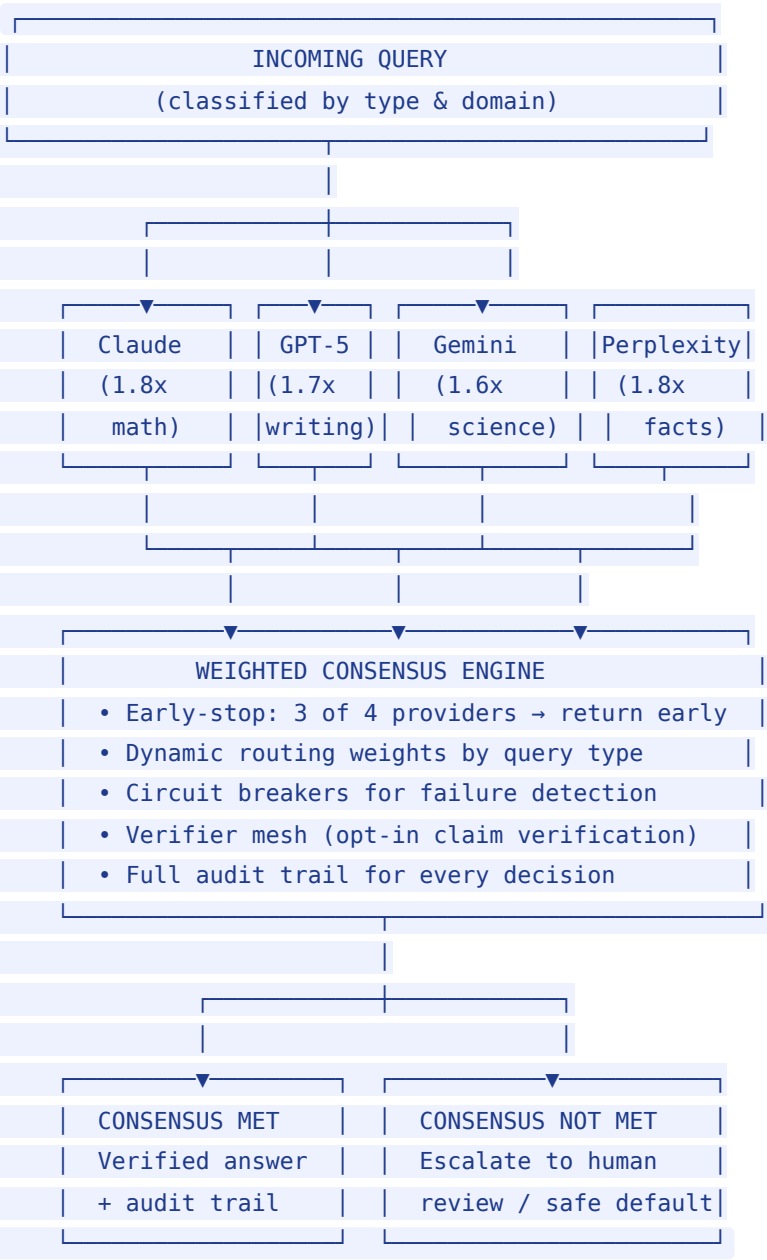
Metric	Value
Total production code	568,934+ lines
Server-side (TypeScript)	~52,600 lines across 100 files
Client-side (React/TypeScript)	~32,000 lines across 846 files
Student Portal (HTML/CSS/JS)	36,587 lines (single self-contained file)
Parent Portal (HTML/CSS/JS)	11,637 lines (single self-contained file)
Language	TypeScript / Node.js
Frontend framework	React 18, Tailwind CSS, shadcn/ui
Database	PostgreSQL with Drizzle ORM
Deployment	Cloud-agnostic (currently deployed on Replit, portable to AWS/Azure/GCP)

KEY SERVER MODULES

Module	File	Purpose
Quad-AI Orchestrator	<code>quad-ai-orchestrator.ts</code> (936 lines)	Core consensus engine — parallel execution, weighted voting, early consensus
Consensus Pipeline	<code>consensus-pipeline.ts</code>	Verifier mesh — claim extraction, cross-referencing, confidence decomposition
Circuit Breaker	<code>ai-circuit-breaker.ts</code> (299 lines)	Per-provider fault tolerance with automatic recovery
Enhanced Cache	<code>cache/ai-cache.ts</code> (629 lines)	Redis + in-memory hybrid cache with cost tracking
Health Monitor	<code>ai-health-monitor.ts</code> (302 lines)	Real-time provider health checking with cached status
Health Endpoint	<code>ai-health-endpoint.ts</code> (305 lines)	Production health API with comprehensive diagnostics
Structured Logger	<code>utils/logger.ts</code> (274 lines)	Enterprise JSON logging with correlation IDs
Request Validation	<code>ai-validation.ts</code>	Zod schema validation on all inputs

Module	File	Purpose
Response Validation	quad-ai-validation.ts	Output contract verification
Streaming	ai-streaming.ts	SSE streaming for OpenAI and Anthropic
Provider Modules	anthropic.ts , gemini.ts , perplexity.ts , openai.ts	Individual provider adapters

Technical Architecture



Three-Portal Ecosystem

Portal	Innovations	Audience	Privacy
Student Portal	83 AI	K-12 students	Completely private — no employer or parent data

Portal	Innovations	Audience	Privacy
	features		exposure
Parent Portal	13 AI features	Parents	Private — employer never sees family data
Enterprise Portal	22 AI features	HR / Benefits teams	Aggregate metrics only — no individual student data

Privacy Architecture: Enterprise sees ONLY aggregate metrics. Parent and student data entirely private. Session data ephemeral. FERPA/COPPA compliant by design, not bolted on.

Cross-Industry Applications

The consensus engine is provider-agnostic and industry-agnostic. Same architecture, different routing weights:

Industry	Application	Consensus Value
Healthcare	Clinical decision support, drug interaction alerts	Digital tumor board — 3 of 4 models must agree before recommendation reaches clinician
Legal	Contract review, case law verification	Four-eyes principle with reproducible decision logs
Financial Services	Risk scoring, KYC, trade pre-clearance	Model committee with regulator-ready consensus trails
HR / Employee Benefits	AI-powered tutoring as premium benefit	Privacy-first ROI metrics without surveillance
Cybersecurity	Threat assessment, anomaly detection	Multi-model consensus reduces false positive rates
Customer Support	Response verification, intelligent routing	Higher accuracy with automatic escalation on uncertainty
Compliance	Regulatory document analysis	Auditable multi-model verification

Competitive Landscape

Solution	Providers	Consensus Voting	Production-Ready	Live Features
Quad-AI (This)	4	Weighted voting + verifier mesh	568K+ lines	100+
ChatGPT Enterprise	1	No	Yes	Chat only
Claude Enterprise	1	No	Yes	Chat only
Khan Academy (Khanmigo)	1 (GPT)	No	Yes	~5
Custom Build	Varies	DIY	18-36 months	Start from zero

Rebuild Analysis

Factor	Estimate
Codebase Reconstruction	Multi-million-dollar engineering effort over 18+ months
Time to Production Parity	18-36 months
Missing Without Production Data	Routing weights, failure pattern corpus, consensus calibration
Domain Expertise Gap	30-year legacy education company's domain knowledge cannot be fast-tracked

As one reviewer noted: *"The real cost isn't the rebuild — it's the opportunity cost of waiting 18 months while competitors ship AI that doesn't hallucinate in healthcare, doesn't contradict itself in legal, and doesn't require a human override in finance."*

Recent Platform Enhancements (March - April 2026)

The platform has been actively hardened and extended in the months immediately preceding this transaction. Buyers performing technical due diligence should be aware of the following recent additions, all of which are live in production:

QUAD-AI ENGINE ENHANCEMENTS

- **Adversarial Verifier Mesh (March 2026)** — Opt-in `verificationMode: 'mesh'` adds claim extraction, support/contradict/omit cross-referencing, independent verification passes, and confidence decomposition (agreementScore, claimCoverage, contradictionPenalty, providerReliability, verificationBonus). Returns synthesized responses structured as Decision / Rationale / Risks / Dissent / Open Questions. This is the acquisition moat — no other consensus system has this layer.
- **Department-Adaptive Routing Weights** — Provider weights adjust per department via `DEPARTMENT_WEIGHT_MODIFIERS`. Legal boosts Claude 1.3x for regulatory precision, Finance boosts GPT 1.2x for quantitative analysis, etc. One API call, six department-tuned outputs.
- **System Status Endpoint** — `GET /api/ai/system-status` provides aggregated system health, provider availability, circuit breaker state, and cache hit rates.
- **Provider Resilience Hardening (April 2026)** — Gemini timeout extended to 45s with single retry; global timeout set to 65s. Eliminates the timeout failures that previously affected long-running queries.

QUAD-AI INDUSTRY TOOLS PAGE ENHANCEMENTS

- **"One API. Every Workflow." Pipeline Visualization** — Live visual showing how a single integration serves Legal, Finance, HR, Operations, Customer Experience, and Marketing.
- **Integration Preview Panel** — Buyers can see real API request/response structure for each department before committing.
- **Dual CTA System** — Distinct paths for technical evaluation (NDA → Architecture Pack) vs. commercial evaluation (white-glove trial).
- **EdTech Integration Block** — Demonstrates LTI 1.3 compliance and direct LMS/SIS integration paths.
- **Social Sharing Layer** — Departmental result sharing for distributed buyer evaluation.

- **Clear Comparison Restoration (April 2026)** — Side-by-side provider comparison view restored and stabilized.

STUDENT PORTAL — INNOVATION SHOWCASE (BUNDLED IN PORTFOLIO 2)

The Student Portal is a single self-contained 37,612-line application. Every AI interaction is verified through the Quad-AI Engine (the "your AI is checked four times before every answer" transparency layer is visible to students on the home screen). The portal is engine-agnostic and can be repointed to any provider stack the acquirer prefers. Live at <https://www.ets-corporate-portal.com/ets-student-portal>. Below are the actual shipped innovations, organized by category.

Quad-AI Powered Learning Tools

- **Quad-AI Neural Challenge Arena** — Multi-AI generated, gradually-harder challenges with consensus scoring
- **AI Learning Time Machine** — Shows the historical context and evolution of any concept the student is learning
- **AI Concept Storyteller** — Converts abstract concepts into narrative form tuned to the student's reading level
- **AI Genius Lab** — Open-ended exploratory environment for deep-dive learning across subjects
- **AI Knowledge Mapper** — Builds visual concept maps showing how the student's knowledge connects across topics
- **AI Study Companion** — Persistent AI study partner the student can return to across sessions
- **Quad-AI World Builder** — Generates immersive learning scenarios for any topic
- **Neural Vision Generator** — On-demand educational image generation tied to the lesson
- **AI Learning Blueprint** (Reality Twin) — Builds a personalized cognitive learning plan from a short profile interview
- **Neural Immersion Arena** — Interactive concept-mastery experiences
- **Expedition Lab** — Guided multi-step learning expeditions across subjects
- **AI SketchPad** — Draw / sketch interface where the AI interprets visual student input

Practical Daily-Use Tools

- **Instant Homework Analysis** — Camera-based scan; Quad-AI extracts the problem, returns step-by-step worked solution
- **Quad-AI Homework Assistant** — Conversational homework help with all four providers verifying every answer
- **Quad-AI Personalized Study Guide** — Auto-generated study guide tailored to the student's profile and recent performance
- **Content Transformer** — Convert any source material into Summary / Step-by-Step / Quiz / Plain-Language Explanation / All-of-the-Above
- **Read Aloud** — Built-in TTS for any content
- **Subject-Specific Help** — One-click Math, Science, English, History tracks
- **Essay Writing Assistant** — Real-time AI suggestions during writing
- **Subject AI Insights** — Per-subject insight panels for student self-reflection

Adaptive Assessment Layer

- **Adaptive Assessment Engine** — Adjusts difficulty in real time based on response patterns
- **Neural AI Test Prep** — AP and standardized test preparation across multiple subjects with adaptive question selection

- **AI-Generated Comprehensive Progress Report** — Multi-dimensional progress report with subject performance, recommendations, and next steps
- **AI-Generated Personalized Assessment** — Custom assessments built from the student's profile

Personalization & Profile

- **Personalized Learning Profile** — Captures learning style, sensory needs, motivation drivers
- **Avatar Creator** — Student-built avatar tied to the profile
- **Adaptive Interface AI** — Accessibility presets; auto-adjusts UI based on detected preferences (font size, contrast, motion, density)
- **Curriculum Generator** — Generates a full curriculum on demand for any subject + grade level

Special Education

- **Neurodivergent Learning AI Suite** — Dedicated tools and supports for neurodivergent learners
- **AI Neurodivergent Learning Intelligence Hub** — Central hub for ADHD, autism-spectrum, dyslexia, anxiety, executive-function, and sensory-processing supports

Live Tutoring (Hybrid Human + AI)

- **Live Tutoring Session** — Request a real human tutor session; AI co-pilot supports the tutor in-session
- **Live Session Joiner** — One-click join for scheduled sessions

Multi-Model AI Showcase

- **Multi-Agent Demo** — Live in-portal demonstration of 50+ AI models running side-by-side, showing the breadth of the underlying engine
- **Neural Interface** — Experimental neural-style interaction surface for next-generation input modalities
- **AI Learning Hub Tour** — Guided onboarding that walks new students through every learning module

Immersive / WebXR Layer (Acquirer-Extendable)

- **A-Frame 1.4.0 + WebXR Integration** — Real WebXR feature detection (`navigator.xr.isSessionSupported('immersive-vr')`) and live `<a-scene>` rendering; works in any connected VR headset
- **enterWebXRMode** — Functional immersive entry point with proper exit and non-VR fallback paths
- **AR/VR Learning Module** — Subject-specific VR learning launchpad
- **Personalized VR Journey + Neural VR Learning Journey** — Modal-driven VR experience surfaces (curriculum content intentionally minimal, designed for the acquirer to extend — explicitly commented in code as *"Ready for Acquirers to Extend"*)

Always-On Safety Layer

- **Quad-AI Four-Check Transparency** — Student-visible badge on every answer showing four-provider consensus verification
- **AI Help Center** — Guided help for using the platform itself

Full feature manifest (every innovation, every code reference) available in the technical data room.

PARENT PORTAL — INNOVATION SHOWCASE (BUNDLED IN PORTFOLIO 2)

The Parent Portal is a self-contained 11,637-line application built around the trust badge "Built for Children. Verified Before Every Answer." Live at <https://www.ets-corporate-portal.com/ets-parent-portal>. Every feature below is verified directly against production code.

Flagship: Neural Mirror™ — Real-Time Cognitive Analysis

- **Neural Mirror™** — Revolutionary cognitive analysis system that monitors the child's actual learning state in real time when hardware is connected (camera, microphone, optional wearable)
- **Simulated Preview Mode** — When hardware is not connected, the same UI runs in clearly-labeled simulation mode so parents can see exactly what Neural Mirror tracks before opting in
- **Neural Mirror™ Cognitive Report** — On-demand multi-dimensional cognitive report generated from live tracking sessions
- **Cognitive Report Badges** — Every metric labeled by `trackingMode` (Live vs. Simulation) so parents always know which numbers reflect real behavior
- **Pre-Confusion Detection Alerts** — Detects the moment a child is about to disengage *before* confusion sets in; surfaces alerts in the Pre-Confusion Detection History panel with one-click "Start Live Tracking" recovery
- **Pre-Confusion Detection Details Modal** — Drill-down for each alert with session context

Report Card Intelligence

- **AI Report Card Narrator** (Feature 5) — Parent uploads the child's actual report card; AI generates a plain-language narrative explaining what the grades really mean for that child
- **AI Report Card Scanner & GPA Tracker** (Feature 5B) — Vision-API-powered extraction of grades from any uploaded report card (image or PDF), with longitudinal GPA tracking across terms
- **Competitive Benchmarking AI** (Feature 6) — Uses uploaded report cards to benchmark the child against state-level peer percentiles; clearly labels simulated example children vs. the parent's real children

Predictive & Forecasting Tools

- **Grade Forecasting** — Per-child forecast cards with "Refresh with AI" generating personalized grade predictions through the Quad-AI engine
- **Success Simulator** — Parent adjusts study hours / intensity and sees AI-predicted outcomes per child

Conversational AI

- **Parent Coach** (AI Academic Guidance Assistant) — In-portal chat for learning strategies, grade-level expectations, study tips, and parenting-for-learning support
- **Quad-AI Learning Insights** — Per-subject insight panels generated by the consensus engine

AI-Powered Action Tools

- **AI Worksheet Generator** — Generate, preview, regenerate, export to PDF, or print custom worksheets matched to the child's grade and gaps
- **Academic Strategy Engine** — Honest, parent-input-driven strategic action plans tailored to the child's actual learning profile (no generic recommendations)
- **Social Story Generator** (Feature 11) — Custom social stories for any scenario (e.g., birthday party, first day at new school) and grade level — autism-spectrum and neurodivergent support
- **IEP Generator** — Multi-disability selection; generates draft IEP language matching how special education plans are actually written

Voice & Tone Analysis

- **Voice Analysis** — Web Audio API integration tracking `stressLevel`, `volume`, `variability`, and `tone` during parent-child conversations

Family & Multi-Child Management

- **Student Profiles** — Add, edit, and switch between multiple children
- **Add Another Student** — Multi-child households fully supported
- **Learning DNA Profile** — Per-child profile capturing learning style, sensory needs, motivation drivers, communication preferences
- **Detailed Progress Report** — Comprehensive academic + cognitive progress over time, print/PDF-ready
- **Subject Performance Analysis** — Drill-down per subject with strengths, gaps, and trend lines

Live Human Tutoring

- **Live Tutoring Request** — Parent submits child's profile through the live tutoring intake; the human tutoring team reviews and pairs the family with the right human tutor (form-based intake, not algorithmic matching)
- **Schedule Consultation** — Booking flow for parent-teacher / parent-tutor sessions

Growth Framework

- **Growth Beyond Grades** — 7-category framework module presenting the dimensions that AI-augmented learners actually develop (Problem Decomposition, Critical AI Literacy, Cross-Subject Transfer, Self-Directed Learning, Confidence Trajectory, Question Quality, Verification Instinct)

Sharing & Account

- **Share Report** — Securely share progress reports with co-parents, teachers, or extended family
- **My Account / Parent Portal Access** — Account management and access controls

Child-Safety & Trust Layer

- **"Built for Children. Verified Before Every Answer."** — Persistent trust badge on every parent-facing screen
- **XSS Hardening** — Markdown rendering pipeline escapes user-supplied characters before conversion; all parent-supplied inputs sanitized
- **COPPA + FERPA Compliant** — Privacy-first; no child data leaves the household account

Full feature manifest available in the technical data room.

RECENT PARENT PORTAL ENHANCEMENTS SUMMARY (APRIL 2026)

- **Voice Analysis** — Web Audio API integration tracks `stressLevel`, `volume`, `variability`, and `tone` during parent-child conversations to surface engagement and emotional patterns.
- **Learning DNA Profile** — Dynamic dropdown for multi-child households with per-child preferences panel covering learning style, sensory needs, motivation drivers, and communication preferences.
- **Academic Strategy Engine** — Generates honest, parent-input-driven strategic action plans tailored to the child's actual learning profile, not generic recommendations.
- **Cognitive Report Badges** — Report modal differentiates metrics by `trackingMode` (Live measurement vs. Simulation) so parents always know which numbers reflect real behavior.
- **Growth Beyond Grades Section** — Informational module presenting the 7 measurement categories that AI-augmented learners actually develop (Problem Decomposition, Critical AI Literacy, Cross-Subject Transfer, Self-Directed Learning, Confidence Trajectory, Question Quality, Verification Instinct).

- **XSS Hardening** — Markdown rendering pipeline now escapes user-supplied characters before conversion and sanitizes all parent-supplied inputs.

IEP GENERATOR ENHANCEMENT

- **Multi-Disability Selection** — IEP generation now supports multiple concurrent disability classifications, matching how special education plans are written in practice.

COMPLIANCE & SECURITY POSTURE

- **Four-Layer K-12 Guardrails** — (1) Specialized pedagogical system prompts, (2) multi-AI consensus verification, (3) input sanitization, (4) domain-locked routing.
- **COPPA + FERPA Compliant** — Privacy-first architecture verified across student, parent, and enterprise portals.

CONSUMER REVENUE LAYER (LIVE)

- **Stripe + PayPal** — Live payment processing for student portal access (\$29.95 one-time) and Report Card Analyzer (\$9.99). Production-ready, not stubbed.

Inside Portfolio 2: AI Growth Measurement Framework

Bundled inside Portfolio 2 alongside the Student & Parent Portals because the Student Portal and the Measurement Framework are functionally connected and can be wired together for multi-student cohort testing and demos. Also included automatically in Portfolio 3 (Full Platform Acquisition).

THE PROBLEM IT SOLVES

Traditional assessment (test scores, grades) measures static knowledge. It cannot measure what AI-augmented learners actually develop: how well they decompose complex problems, how critically they evaluate AI output, how they transfer reasoning across subjects, and how they direct their own learning. As AI becomes ubiquitous in classrooms and corporate L&D, the entire measurement industry needs a new framework — and there isn't one.

BEYOND TEST SCORES — 7 MEASUREMENT DIMENSIONS

1. **Problem Decomposition** — How learners break complex problems into solvable parts
2. **Critical AI Literacy** — Whether learners verify or blindly trust AI output
3. **Cross-Subject Transfer** — Whether reasoning learned in one domain applies to others
4. **Self-Directed Learning** — Whether learners drive their own inquiry
5. **Confidence Trajectory** — How conviction evolves with new evidence
6. **Question Quality** — Whether learners ask better questions over time
7. **Verification Instinct** — Whether learners independently fact-check

LIVE IMPLEMENTATION

Working dashboard at <https://veridect.ai/ai-growth-measurement> — 23 pre-loaded students, classroom-level cohort analysis, per-student radar charts, category scorecards, tool-usage breakdown, and multi-student comparison view. Backend APIs at </api/ai/growth-analysis> and </api/ai/simulate-student> accept any properly formatted student data.

STRATEGIC FIT & CROSS-INDUSTRY APPLICABILITY

- Plugs naturally into LMS, SIS, and assessment platforms

- Vendor-neutral by design — works with any AI tool stack
- Industry-agnostic methodology — adaptable to corporate L&D, professional training, healthcare education, military training
- Significant implementation work already completed: measurement engine, dashboard, cohort visualization, integration APIs

BUYER POOL

Distinct from the Quad-AI Engine buyer pool — natural acquirers include assessment companies (ETS, Pearson, NWEA, Curriculum Associates), LMS platforms (Instructure, PowerSchool, Schoology), corporate L&D platforms (Cornerstone, Workday Learning, Degreed), and AI infrastructure companies seeking a measurement layer for AI-native applications.

The framework dossier is available under separate cover.

Live System Access

Resource	URL
Quad-AI Industry Tools (6 enterprise departments, live verifier mesh, confidence decomposition)	veridect.ai/quad-ai-demos
Student Portal (Portfolio 2 — 83 AI innovations, K-12)	www.ets-corporate-portal.com/ets-student-portal
Parent Portal (Portfolio 2 — Voice Analysis, Learning DNA, Academic Strategy Engine, Growth Beyond Grades)	www.ets-corporate-portal.com/ets-parent-portal
Enterprise Portal (employer/HR-facing analytics layer)	www.ets-corporate-portal.com/ets-enterprise
AI Growth Measurement Dashboard (7-dimension framework, 23-student cohort)	veridect.ai/ai-growth-measurement
Quad-AI Education — Live (real-time provider responses, latency, consensus)	https://edtutorial.services/quad-ai/
AI Sandbox (full platform, live)	https://www.edtutorial.services/ai-sandbox/
Employer Benefits Portal	https://www.edtutorial.services/employers-enhanced/

Documents Available

Document	Purpose	Best For
This Document	Architecture review + full technical due diligence	First share, technical evaluation
Multi-Provider Analysis	10 AI providers analyze 6 industries (60/60 analyses)	Industry-specific value
Platform Overview	Full innovation inventory, pricing structure	Comprehensive scope review
Tech Sales Package	Technical deep-dive for engineering teams	CTO/engineering evaluation
Sale Package	Deal structure and pricing	Late-stage negotiation

Contact: Lisa Russell / CEO President, Educational Tutorial Services

lrussell@educationaltutorialservices.com

Legacy education company (est. 30 years) | Half a million children served | Fortune 100 partnerships

This document contains proprietary information. Distribution limited to authorized recipients under NDA.