

VerdictTank v3.6 - Critical Review

4-Judge Pipeline Verdict - August 11, 2026 - Review ID: VT-36-2026-001

CONDITIONAL GO - 3-1 Majority

Reviewed: [VerdictTank v3.6 Business Proposal](#) + [Technical Architecture](#) + [Minimum Viable Legal Framework](#). This v3.6 revision was submitted in response to a prior 4.1/10 Conditional Go with 9 conditions. All 9 conditions are addressed.

Meta: v3.6 is a post-critical-review rewrite. All 9 conditions from the v3.5 review are present with dedicated sections. This review evaluates whether the substance behind those sections justifies a Go verdict.

Dissenting Opinion - Judge 3: "Go (Unconditional). The team has systematically de-risked the project. This is no longer a speculative bet; it is a well-structured plan to validate and build. The correct action now is to execute that plan."

Verdict Summary

JUDGE	VERDICT	COMPOSITE
Judge 1	Conditional Go	4.8/10
Judge 2	Conditional Go	6.0/10
Judge 3	GO	-
Judge 4	Conditional Go	-

MAJORITY RULING

3-1 (75% consensus)

10-Dimension Scores

#	DIMENSION	JUDGE 1	JUDGE 2	CONSENSUS
1	Problem Definition	6	-	6.0
2	Market Analysis	3	-	3.0
3	Competitive Moat	4	4	4.0
4	Business Model	5	-	5.0
5	Unit Economics	6	-	6.0
6	Technical Architecture	7	-	7.0
7	Go-to-Market Strategy	3	-	3.0
8	Risk Assessment	7	-	7.0
9	Execution Feasibility	5	3	4.0
10	Growth Trajectory	2	-	2.0
COMPOSITE		4.7/10		

What improved (was 4.1): Architecture (7), Risk Assessment (7), Unit Economics (6). Legal framework, labeling discipline, and downside modeling are real progress. **What didn't:** Growth (2), Market (3), GTM (3) were relabeled, not re-derived.

Unanimous Flaws (All 4 Judges Agree)

#1: Zero user-validated signal. Every capability sourced from Competitive Landscape Research Batch 001. The proposal is honest

(Section 02a: "None of the nine capabilities carry user-validated status") but honesty does not substitute for evidence. The beta survey hasn't run.

#2: No acquisition strategy. Forecasts 210/780/2,450 paying users by Year 3 with zero acquisition channel, CAC assumption, paid/organic mix, or top-of-funnel mechanism. "SEO" appears nowhere. The growth forecast is a conversion table without a distribution story.

#3: Growth trajectory unchanged from v3.5. Same 3-year hockey-stick, just relabeled. Year 1-3 numbers were raised over v3.5 using mechanisms that are either deferred (White-Label) or un-surveyed (Chat-to-Refine per Section 02a).

#4: White-Label priced with deferred isolation. Architecture defers RLS/multi-tenant isolation (1.1, 1.2) while pricing White-Label at \$1,999+/mo (Section 11). One onboarded partner without RLS = injunctive risk day one.

Novel Insights (Not in v3.5 Review)

JUDGE 2 Fix-It scope contradiction between roadmap and architecture. Roadmap (Section 12) calls Fix-It "single-tier initial." Architecture (Section 5.5) specs the full 3-tier quality rubric evaluator with an additional per-action-item judge-model call. Both documents can't be right.

JUDGE 2 Dual-score validation gate is theater. The pass threshold of "majority preference" from a self-selected beta cohort guarantees a pass. Reject requires a 65%+ threshold or an equally-salient single-score option. The gate is calibrated for confirmation, not validation.

JUDGE 2 No GTM contingency for dual-scoring failure. Dual scoring is the stated moat. The architecture can revert, but the marketing plan has no answer for "what's our differentiator if dual scoring fails validation?"

JUDGE 2 MVP timeline is longer than advertised. Beta survey not counted in 8-week window. If survey takes 2-3 weeks, real timeline is 10-11 weeks. Both documents present "8 weeks" as committed.

JUDGE 4 URL-to-Review creates third-party PII risk. Snapshots of user-submitted URLs ingest third-party PII with no legal basis operationalized. One takedown request triggers a privacy incident.

JUDGE 3 MVL creates external counsel dependency. White-Label gated on outside counsel review. Timeline is no longer fully within the team's control. Needs named counsel, retainer, and estimated turnaround.

What Survived (Still True)

Competitive moat - empty category. No direct competitor does all 10 dimensions with multi-judge consensus. The full coach-judge-explain-fix-re-score loop is genuinely differentiated.

Unit economics are honest now. \$0.98/review fully loaded, 75-79% margins by tier, churn/cost sensitivity tables. This section went from weak to strong.

Architecture discipline is real. JSONB overlay with validation gate, rollback path, numeric threshold. MVP cut list matches prior Week 1 Cut List component-for-component.

Three Critical Conditions for Go

Condition 1: Re-baseline the timeline. The 8-week MVP ignores the beta survey window and Fix-It scope contradiction. Corrected timeline for one person: 12 weeks (including survey). Publish it.

Condition 2: Remove White-Label from pricing. Replace the \$1,999+/mo tier with "Available post-MVP. Join waitlist." Hard provisioning gate in code that blocks White-Label onboarding until MVL checklist is fully green.

Condition 3: Document dual-score failure contingency. One paragraph in Section 13 acknowledging the possibility and stating the fallback differentiator if the 25-user gate rejects dual scoring.

Priority-Ranked Action Plan

PRIORITY	ACTION	EFFORT
P0	Re-baseline timeline to 12 weeks with survey window shown	30 min
P0	Remove White-Label from public pricing; add waitlist placeholder	10 min
P1	Fix validation instrument: 65%+ threshold, equally-salient single-score option	1 hour

P1	Add GTM contingency paragraph for dual-score failure to Section 13	30 min
P1	Resolve Fix-It scope contradiction (single-tier vs 3-tier)	30 min
P1	Add TAM/SAM/SOM methodology to proposal	3 hours
P2	Run beta survey before committing schema work	Schedule
P2	Engage outside counsel for MVL review	Admin

Reconciled Timeline

PHASE	CLAIMED	CORRECTED	ACTIVITY
Beta Survey	Not counted	Weeks 0-3	Survey Inner Circle cohort, analyze results
Corpus Schema	Weeks 1-3	Weeks 4-6	Schema migration, sanitization gate, 200-sample QA
Dual Scoring	Weeks 4-6	Weeks 7-9	JSONB overlay, GIN index, validation instrument
Fix-It MVP	Weeks 4-7	Weeks 8-11	Single-tier rubric, before/after diff, re-score integration
Chat-to-Refine	Weeks 6-8	Weeks 10-12	REST API, ghostwriting guard, static fallback bank
Total	8 weeks	12 weeks	One person, including survey period

Disclaimer: This review is an automated critique generated by AI models, not professional advice. Scores and verdicts are opinions. Do not rely on this for

investment, procurement, or funding decisions. VerdictTank and its reviewers do not provide legal, financial, or business advice.

VerdictTank Pipeline - Four models. Three architectures. One verdict.
Review conducted August 11, 2026. AI costs: ~\$0.71 total. Pipeline duration: ~5 min.

Proposal (v3.6) - Prior (v3.5) - Architecture - Legal