

JZAI Consensus Tool: Consensus The Truth Engine

User Guide, Architecture
& AI Model Reference

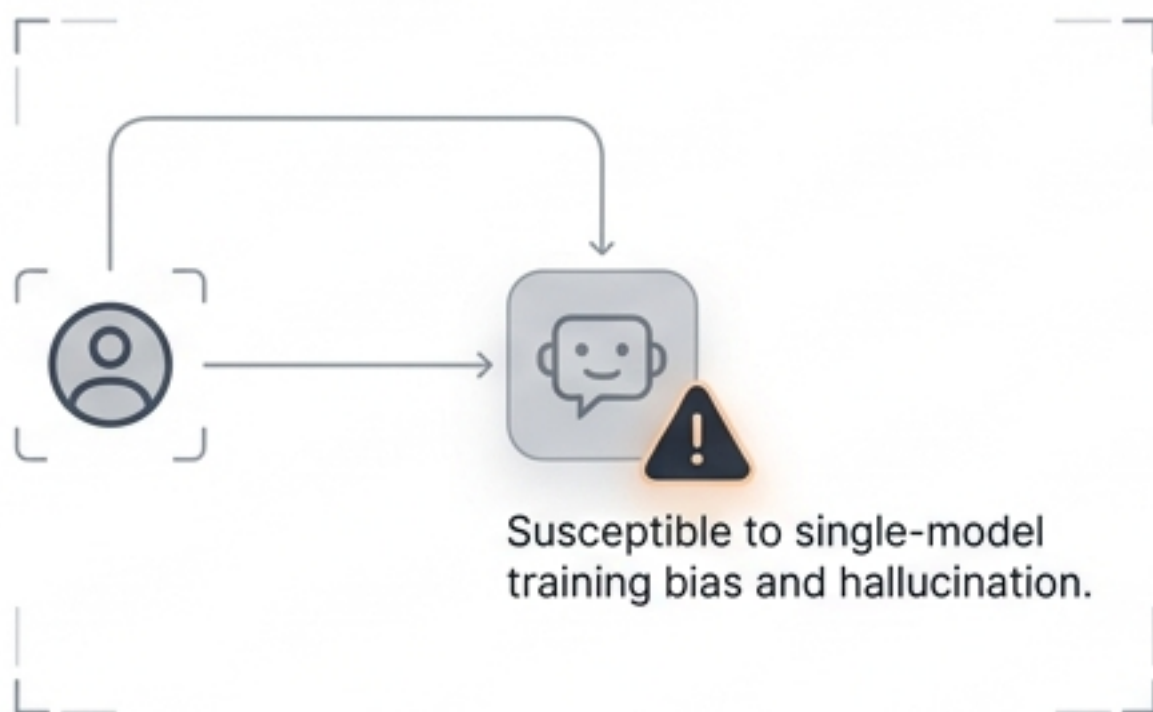
[Version S9]

[August 2026]

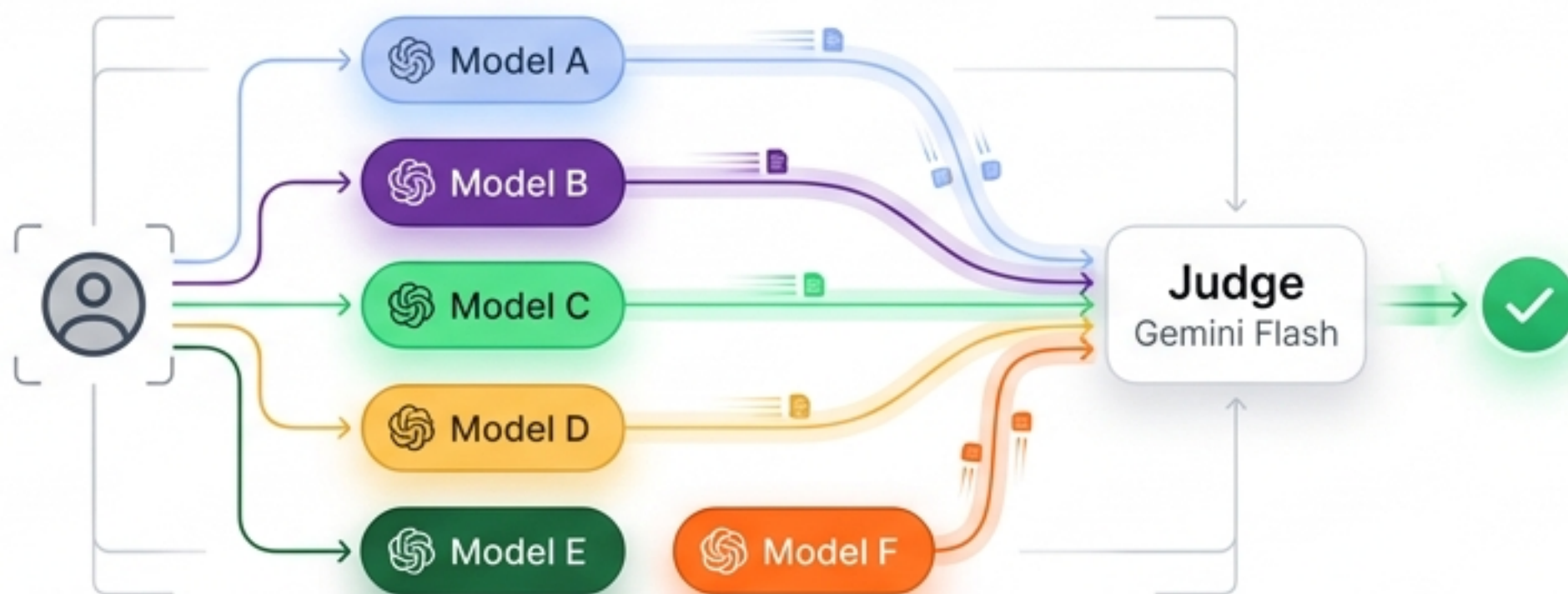
[ai.johnzur.com]

Agreement across multiple independent AIs is more trustworthy than any single AI's response alone.

Legacy: Single-Model Vulnerability



The Consensus Era



Bias Mitigation



Cross-model averaging eliminates the dominance of any single AI's inherent training bias.

Disagreement Visibility



Explicitly calling out where models conflict surfaces genuine uncertainty.

Web Search Grounding



Integrating live web queries ensures real-time factual accuracy.

JZAI Ask Tab

Status Indicator

N AIs ready. Confirms OpenRouter reachability.
The tool functions even if some models are down.

Live Balance Pill

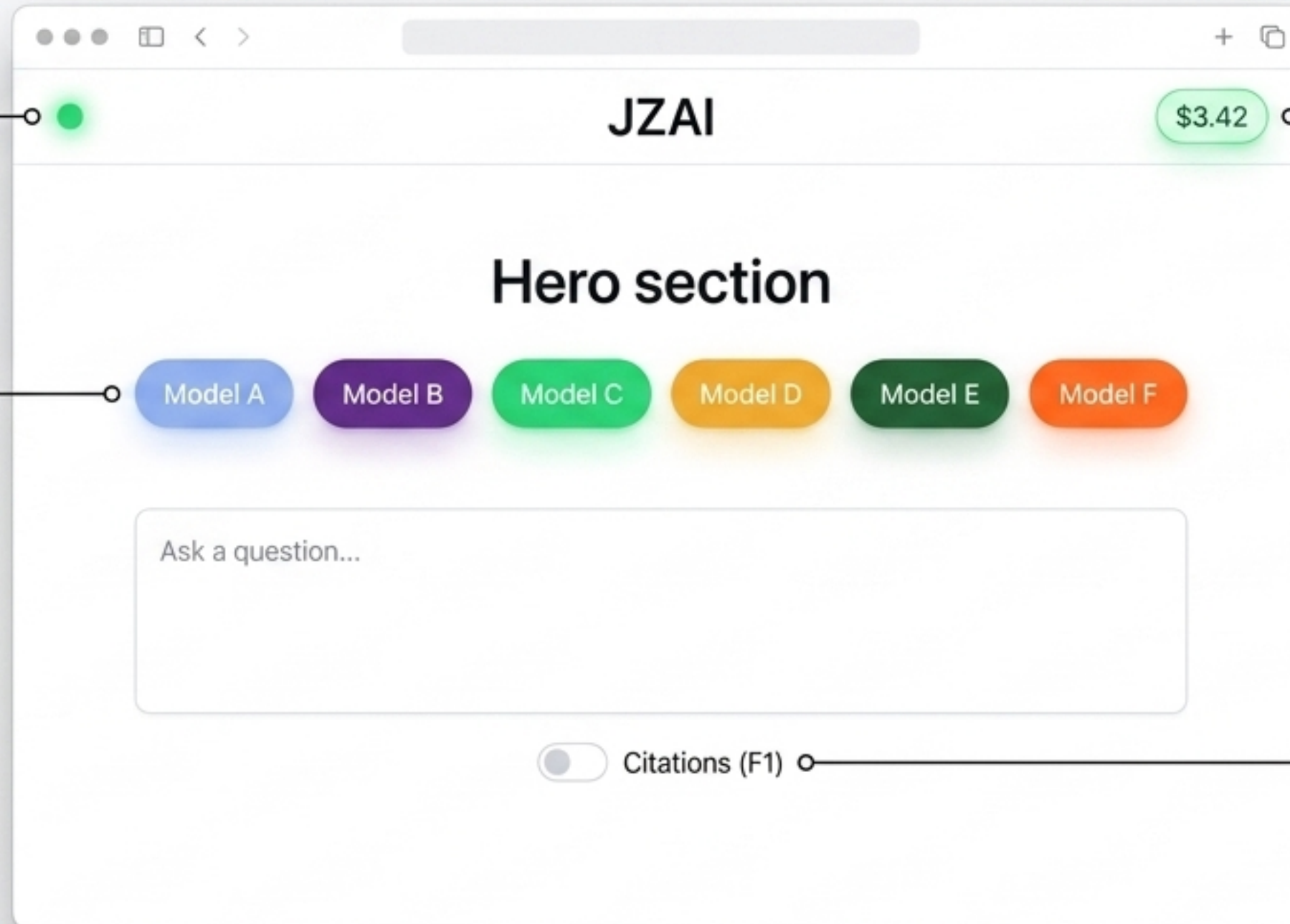
Fetches OpenRouter balance locally.
Green (>\$2.00) → Amber (\$1.00-\$2.00) → Red (<\$1.00).

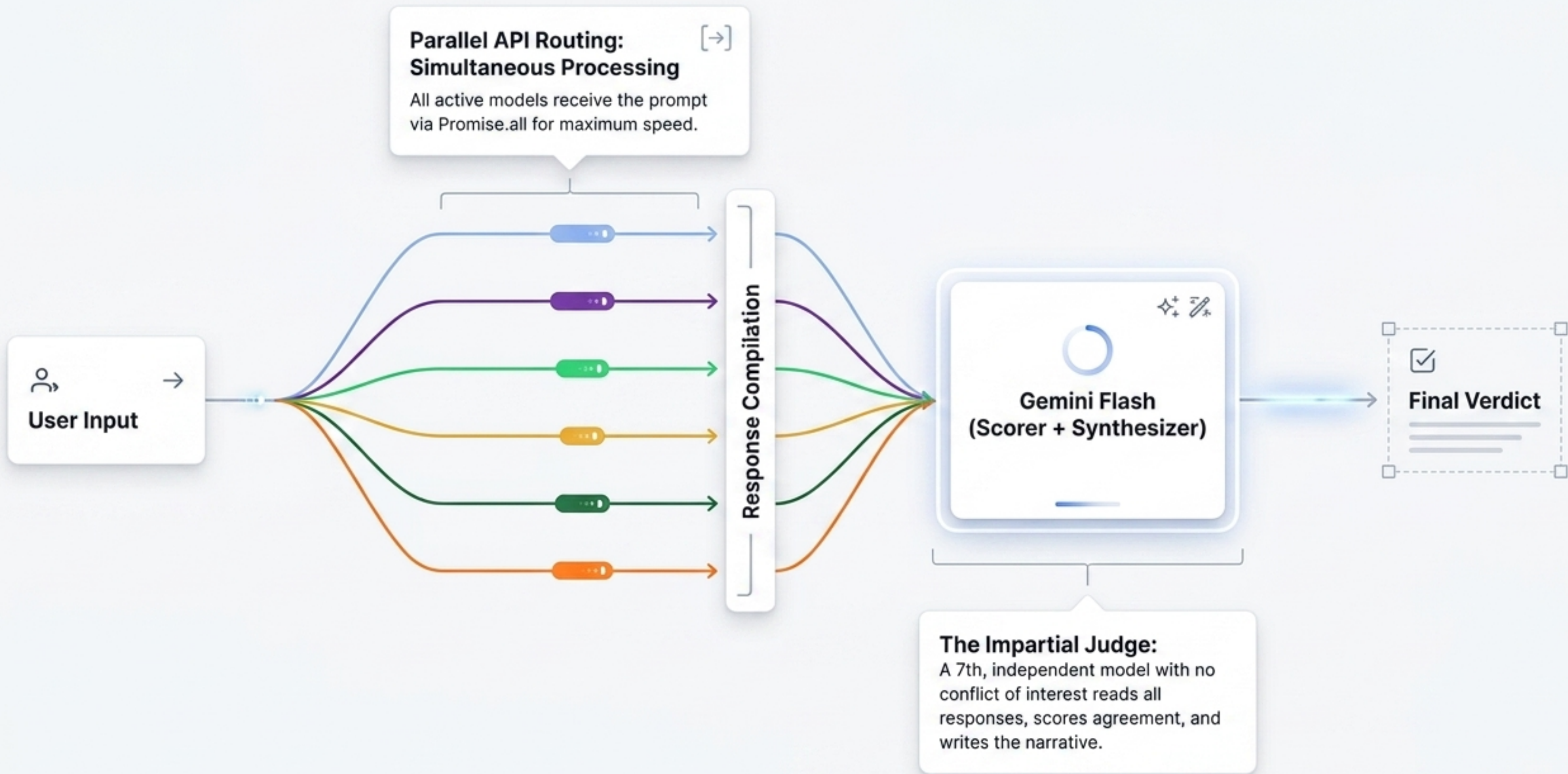
AI Model Pills

Click to toggle AIs on/off.
Settings persist across sessions.
(Minimum 1 active required).

Citations Toggle (F1)

When enabled, appends a citation instruction.
Models will link to specific publications and websites.





Consensus Verdict Anatomy

 Agreement Score: 85

Consensus Verdict

F3 Outlier Detection:

Surfaces unique, highly-valuable insights missed by the majority.

Model A

Indicates positive outlook with caution.

Model B

Strongly suggests negative impact.

Model C

Agrees with positive trend.

Model D

Neutral stance, data inconclusive.

Model E

Supports positive analysis.

Model F

Highlights significant outlier risk.

Key Points

- **Shared Findings:** Models A, C, E agree on core positive drivers.
- **Absolute Agreements:** All models acknowledge data stability.
- Another finding placeholder text.

Disagreements

Model B (Deep Purple) diverges significantly, arguing for negative impact due to [specific reason]. Model F (Vibrant Orange) identifies a unique risk factor ignored by others, suggesting potential volatility. This section details which AI took which position and why.

Bottom Line: Proceed with caution, considering the outlier risk highlighted by Model F.

[1] The Color-Coded Score:

The 0-100 agreement metric.

[2] Per-AI Conclusions:

One-line summaries (8-12 words) of each model's explicit stance.

[3] Key Points:

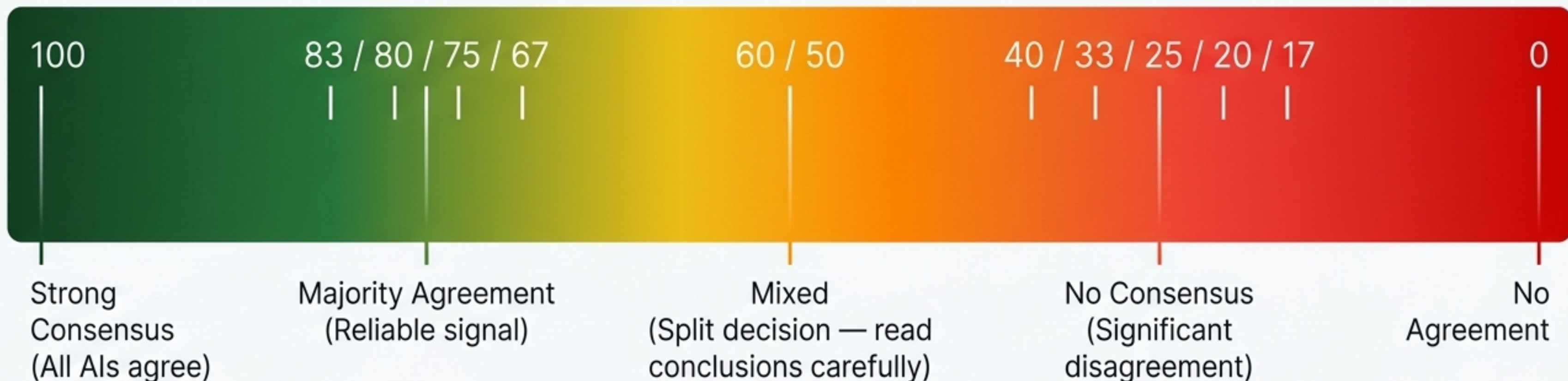
Bulleted list of shared findings and absolute agreements.

[4] Disagreements:

The highest value section—explicitly detailing which AI took which position and why.

[5] Bottom Line:

A direct, one-sentence recommendation.



Local Computation

Score is computed entirely locally in JavaScript:
 $(\text{agreement_count} / \text{total_AIs}) * 100$. Eliminates scorer drift.

Semantic Similarity

Reflects semantic similarity of conclusions, not just a keyword vote count.

F12 Override Flag

Fires “No consensus — verify independently” at $\leq 33\%$ when 3+ AIs are active.



Devil's Advocate (F4)

Fires an adversarial prompt to all active models, explicitly challenging the consensus verdict to test its structural integrity.



What Would Change This? (F8)

A reversal-evidence prompt. Asks each AI what specific future evidence or circumstances would force them to alter their conclusion.



Prior vs. Current Score (F11)

Tracks shifting agreement (e.g., '83% → 60%') in the UI header as conversation deepens on follow-up rounds.



Follow-Up Verdict Comparison (F9)

Synthesizer explicitly states whether new information confirms, weakens, or contradicts the prior round's consensus.



Medical / Clinical

“Should I be concerned about an LDL of 142?”

Value Prop

Disagreement flags genuine clinical uncertainty; Perplexity Sonar pulls live guideline updates.



Financial Decisions

“Move \$60k from a 4.3% HYSA to equities?”

Value Prop

Exposes distinct risk-calibration biases across different models.



Technical / Architecture

“Build my next web app in React or Vue?”

Value Prop

Surfaces defensible tradeoffs rather than a single biased recommendation.

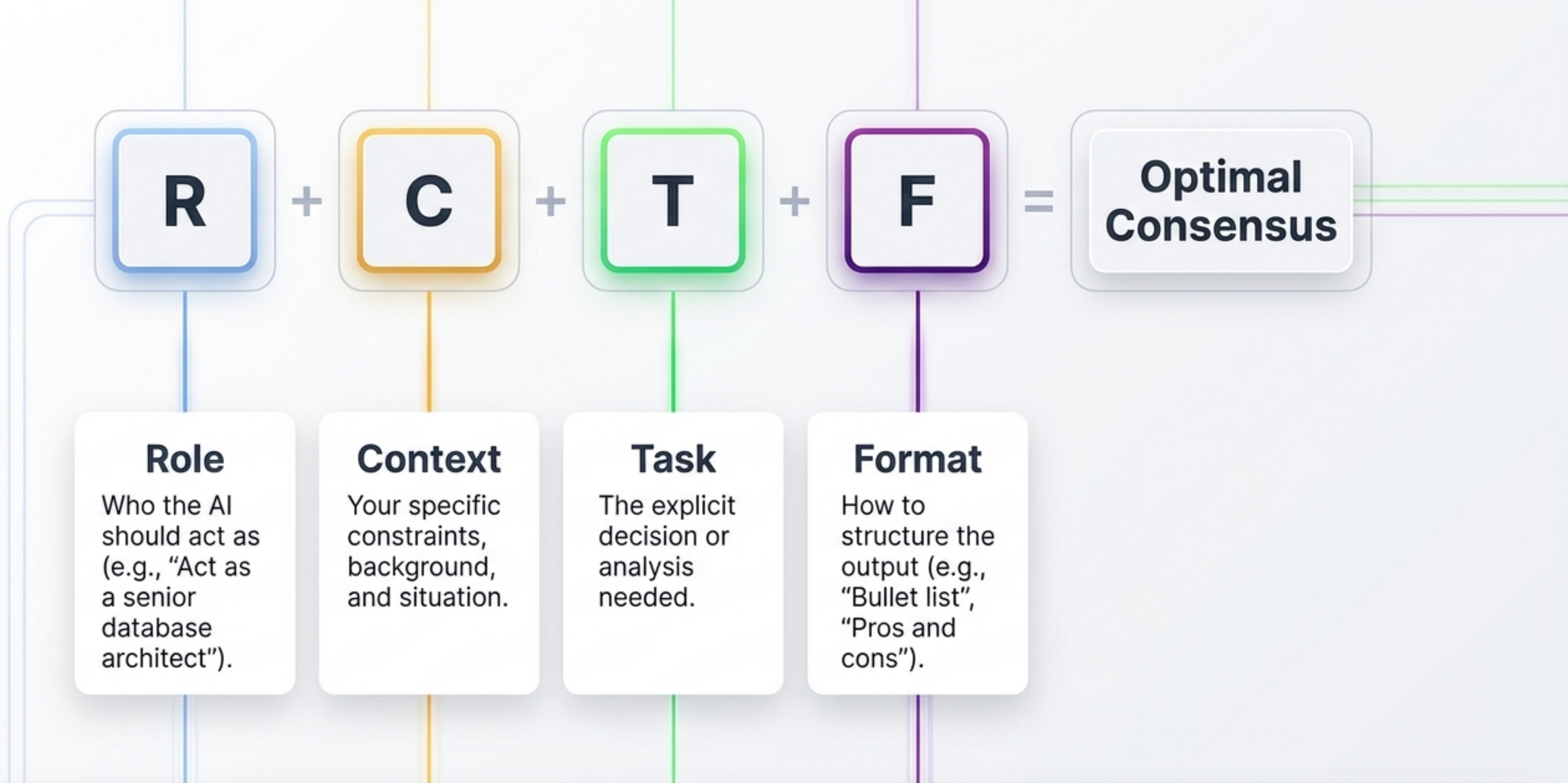


Legal / Contractual

“Differences between an NDA and a non-compete?”

Value Prop

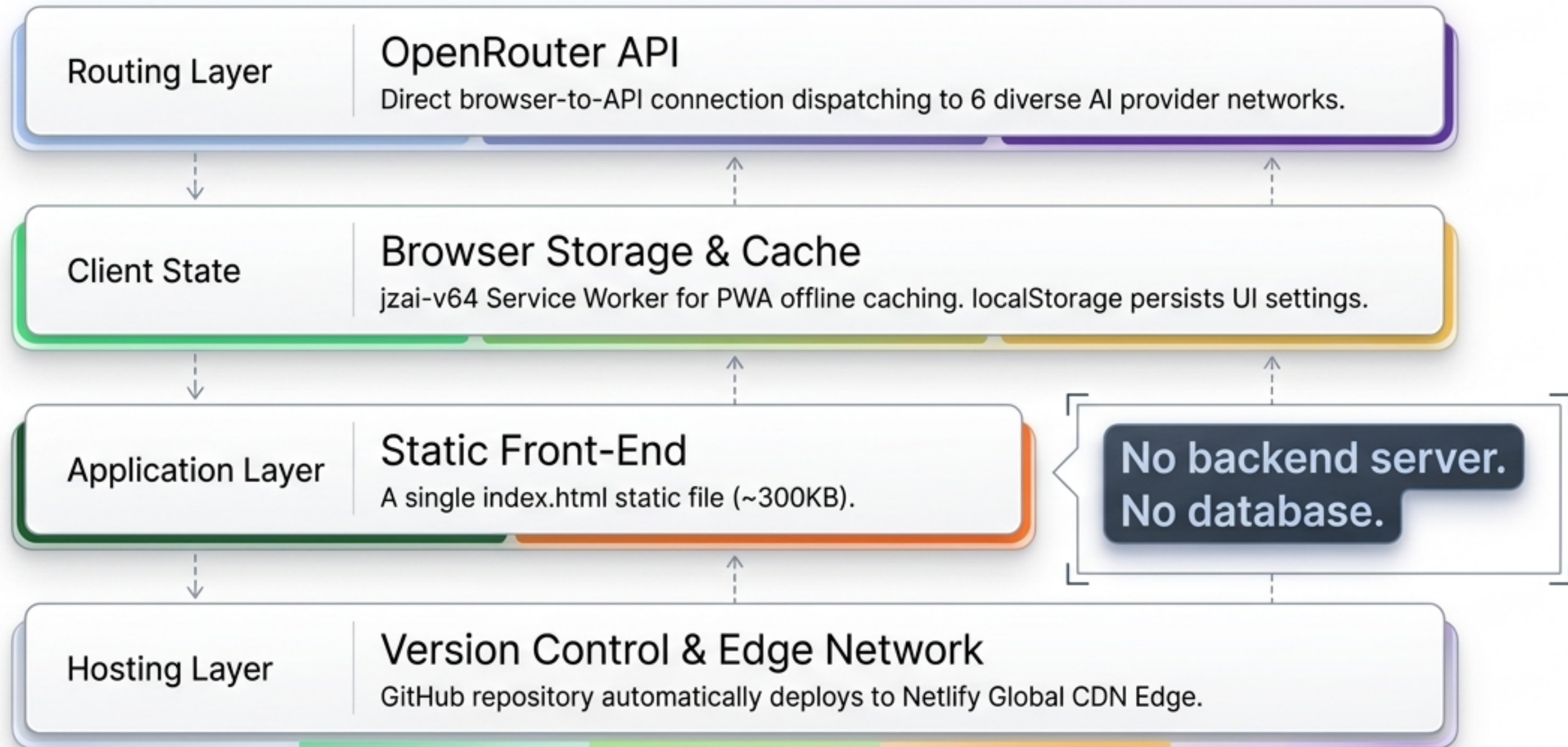
Strong signal on concept accuracy prior to actual legal consultation.



The quality of your initial question directly determines the structural integrity of the final consensus.

Model	Model ID	Unique Trait	Output Token Budget
ChatGPT	openai/gpt-4o	Multimodal flagship	1,000
Claude	anthropic/claude-sonnet-5	Nuanced, careful reasoning	1,500
Gemini 2.5 Pro	google/gemini-2.5-pro	'Thinking model' (Drives cost)	4,000
Meta AI	meta-llama/llama-4-scout	Open-weight, cheapest	1,500
Grok 4.20	x-ai/grok-4.20	Real-time X data, direct	3,000
Perplexity Sonar	perplexity/sonar	Live web search, cites sources	1,000
THE JUDGE: Gemini 2.5 Flash	google/gemini-2.5-flash	The Impartial Scorer & Synthesizer	300 + 500

Serverless Stack



Cost Analysis & Alternative Routings

Digital Receipt		
Worst-Case Token Cost		
	ChatGPT (1,000-token output) 1,000 tokens	\$0.00800
	Gemini 2.5 Pro (4,000-token output) 4,000 tokens	\$0.06000
	Claude (1,500-token output) 1,500 tokens	\$0.00450
	Grok 4.20 (3,000-token output) 3,000 tokens	\$0.00300
	Perplexity Sonar (1,000-token output) 1,000 tokens	\$0.00200
	Meta AI (1,500-token output) 1,500 tokens	\$0.00120
	Scorer & Synthesizer (500-token output) 500 tokens	\$0.00443
Total Cost: \$0.08313		
All 6 AIs + Scorer + Synthesizer		

Massive output budget drives the bulk of this cost.

Alternative Routing Configurations

5-Model Config
(No Gemini Pro)

~\$0.04
per query

Top 3 Triumvirate
(ChatGPT, Claude, Perplexity)

~\$0.03
per query

Ultra-Lean
(Meta AI only)

~\$0.004
per query



**The ultimate
command center
for high-stakes
verification.
Trust through
consensus.**

Diagnostics & Telemetry

Symptom

Diagnosis

Status shows
< 6 AIs.

OpenRouter check failed;
tool automatically adapts
and still works.

Cards spin
indefinitely.

OpenRouter balance
depleted OR 60-second
long-query timeout.

Settings don't
persist.

Browser is in
Private/Incognito mode
(localStorage disabled).

Interface won't
update.

Service worker (jzai-v64)
caching old version;
execute a hard refresh.