

High-Dimensional Sequence Predictor

Bayesian Orchestration Control

Position:

**Agentic AI systems should be making
Bayes-consistent decisions!**

<https://icml.cc/virtual/2026/poster/67041>



ICML 2026

The 43rd International Conference
on Machine Learning



JULY 6–11, 2026



SEOUL, KOREA



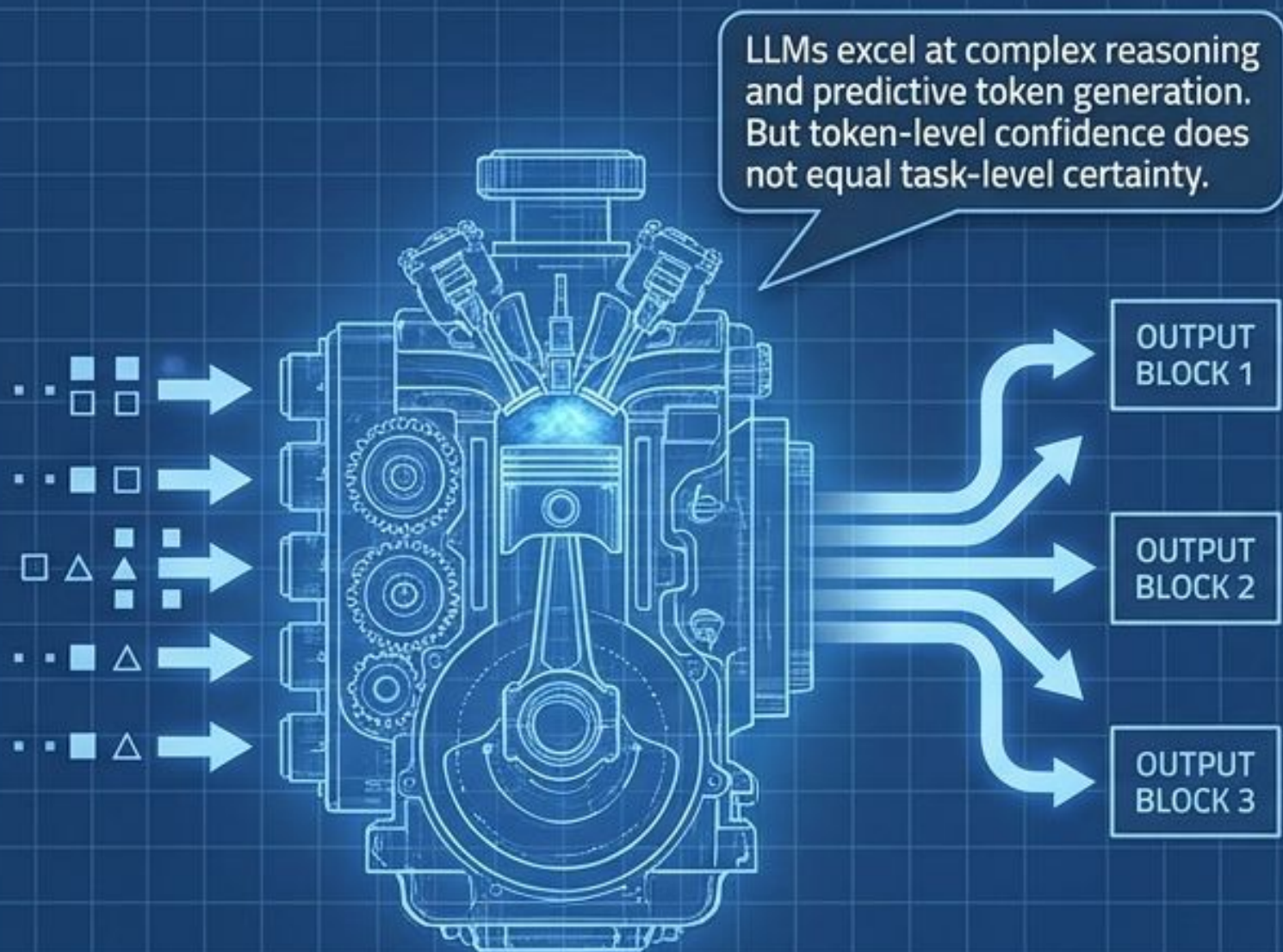
icml.cc/2026

Bayes-consistent
AI systems.

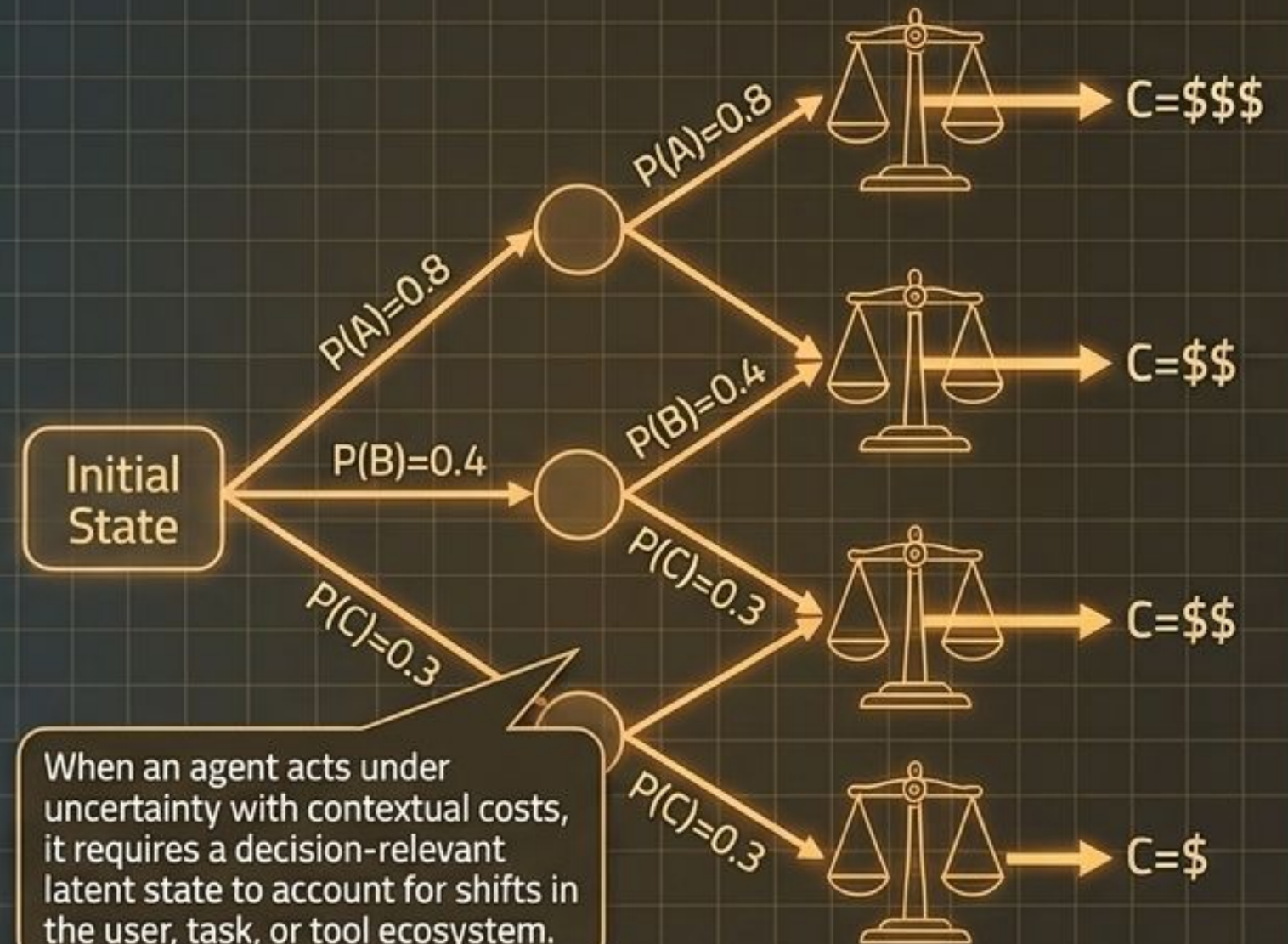
Inspired by the 2025 workshop at MBZUAI
"Rethinking the Role of Bayesianism in the Age of Modern AI"

NotebookLM

Predictive Execution

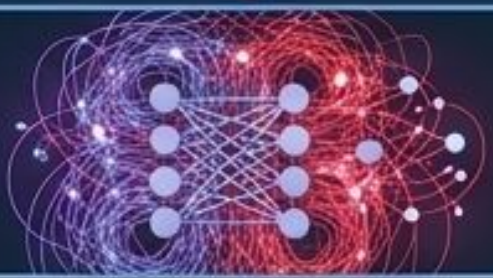
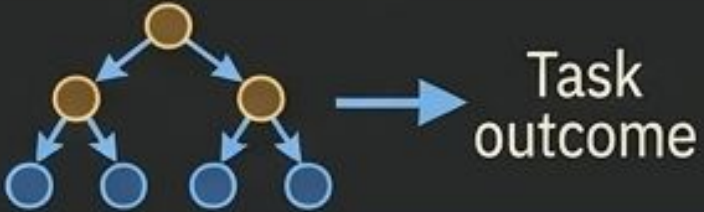
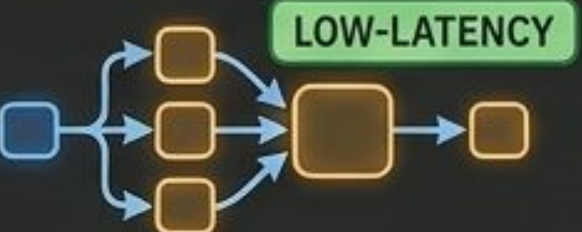


Decision Under Uncertainty



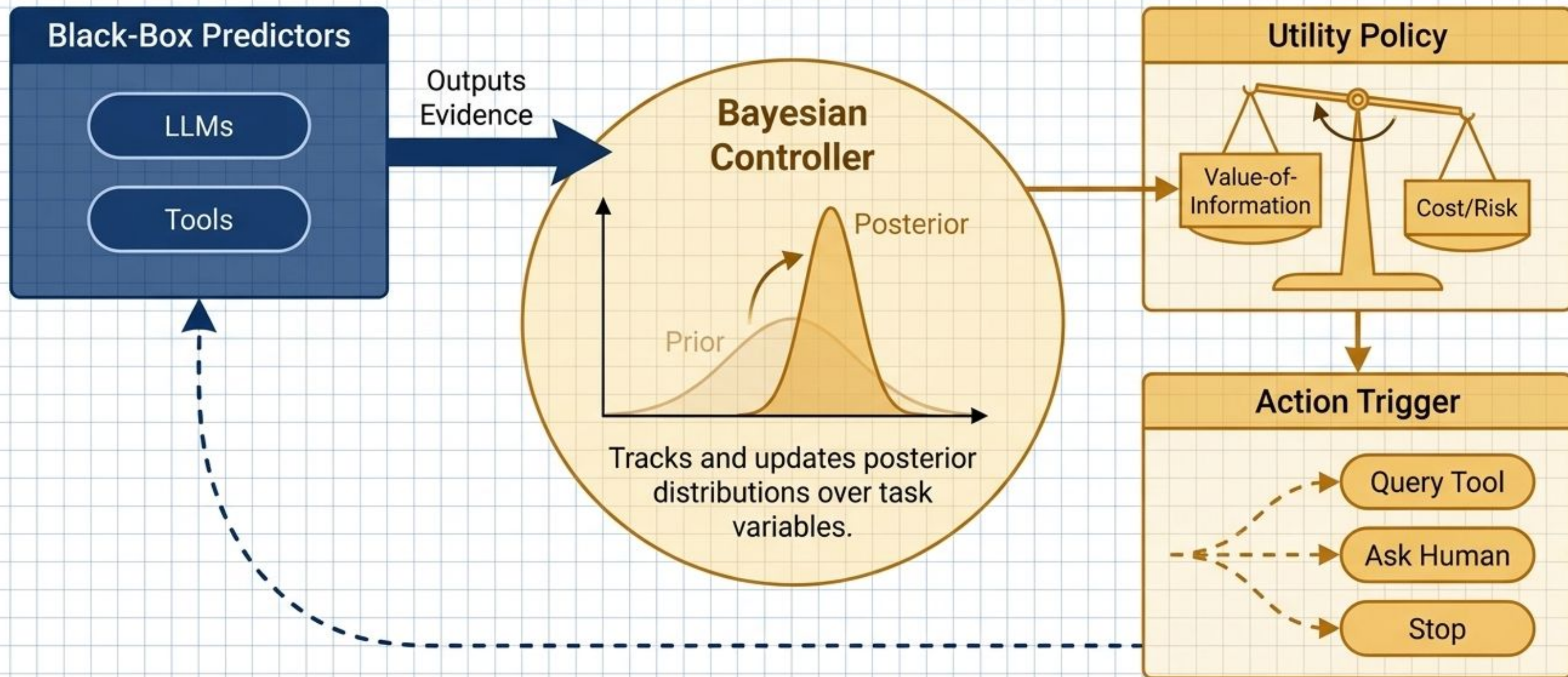
Prediction is about generating the next plausible step.
Orchestration is about routing, stopping, escalating, and allocating budget under risk.

The control layer is where Bayesian principles belong.

	✗ BAYESIAN LLMs	✓ BAYESIAN ORCHESTRATION
Target of Update	Internal LLM Parameters 	Control Layer Latent Variables 
Uncertainty Handled	Syntactic (Token-level)  Token distribution	Semantic (Task/Decision-level) 
Role of the LLM	The solitary decision-maker  Uncertain command	A predictive evidence generator 
Computational Overhead	Prohibitive & Unbounded 	Feasible & Low-Latency 

Takeaway: Treat LLMs as black-box predictive components. Their outputs are probabilistic evidence fed into an independent Bayesian controller.

The **Bayesian Control Loop** separates evidence from decision making.



Triggers tools only when expected improvement outweighs the cost.

Every agent output acts as a **noisy sensor reading**.

$$r_t(y) \propto r_{t-1}(y) \cdot p_{i_t}(z_t | y)^{\alpha_{i_t}}$$

$r_{t-1}(y)$: The Prior Belief



The orchestrator's **confidence** in the task outcome before querying the agent.

$p_{i_t}(z_t | y)$: The Evidence



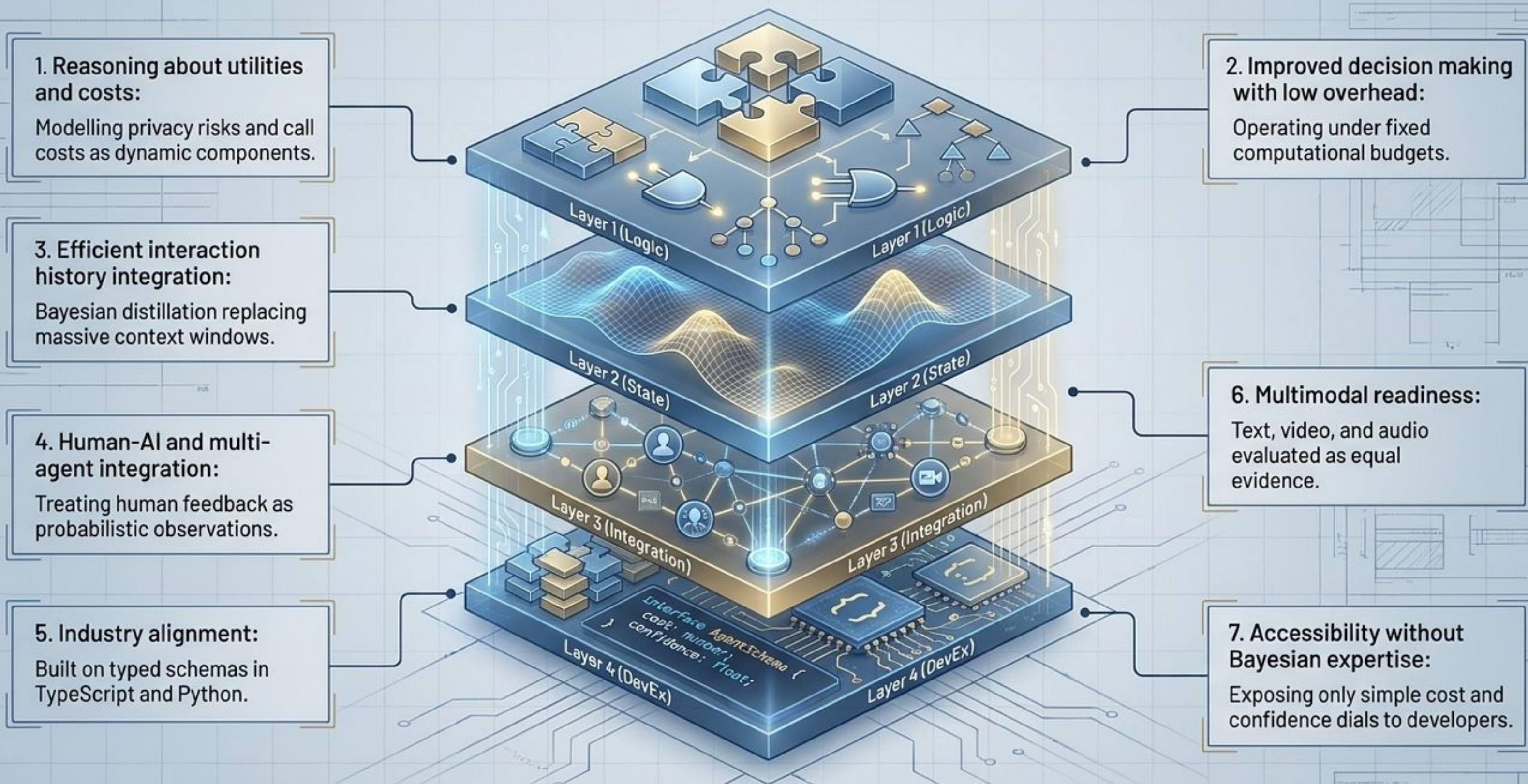
The **observation model** mapping the raw agent message (z_t) to a **likelihood of success**.

α_{i_t} : The Reliability Filter

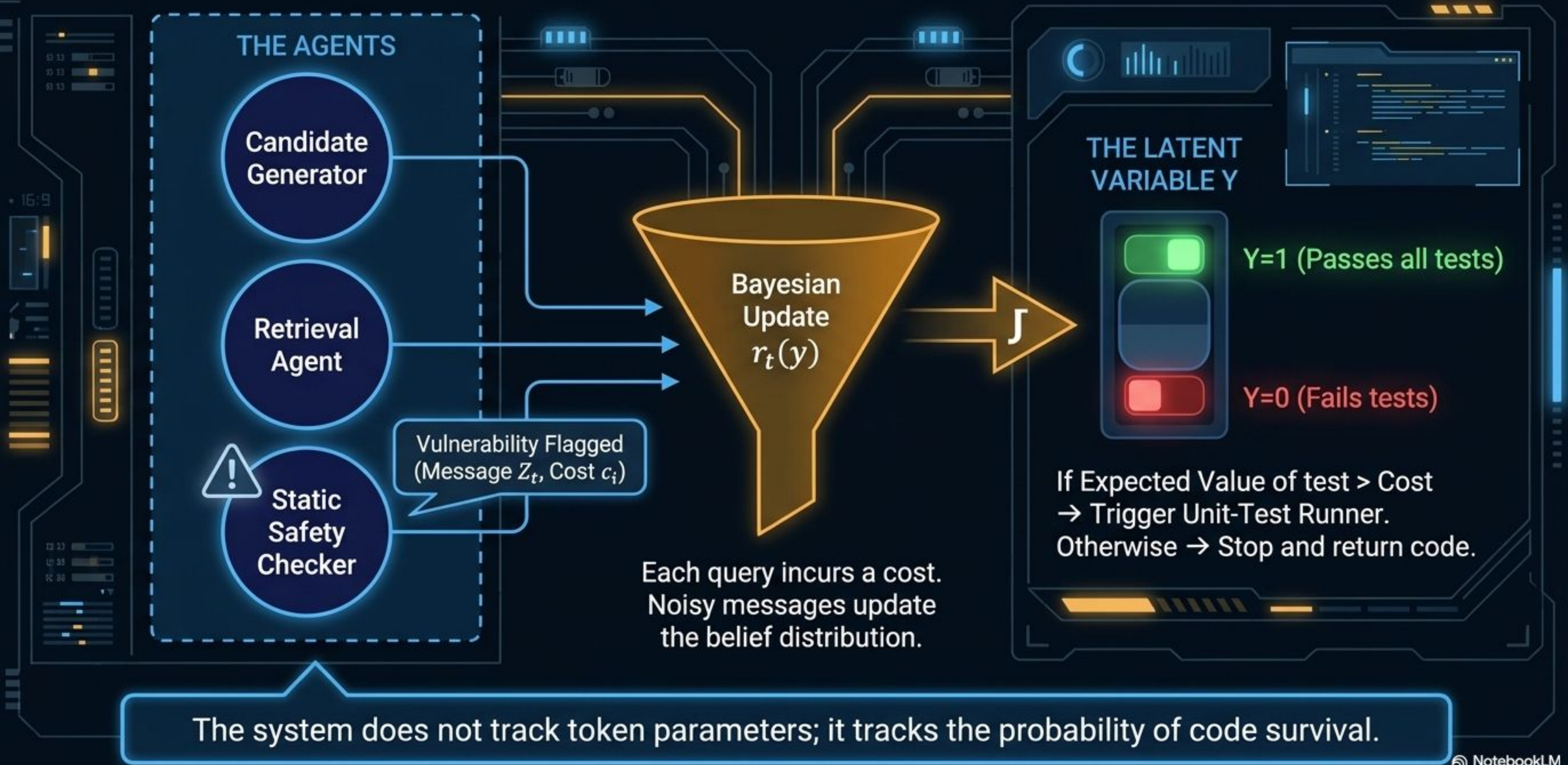


A crucial source-specific weight. Down-weights evidence from weaker or biased sources so the system does not produce overconfident updates.

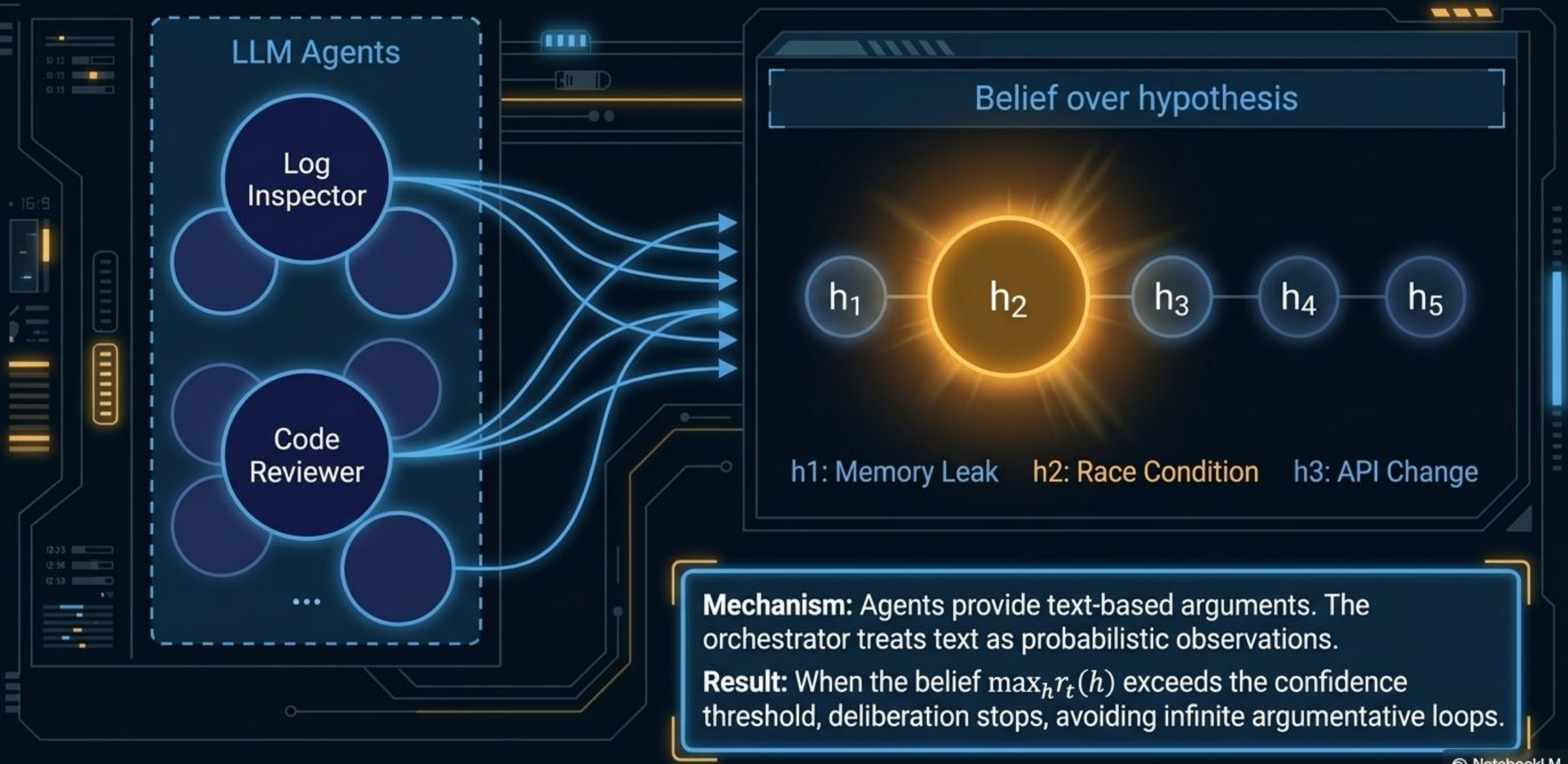
An architecture built for modern human-AI software stacks.



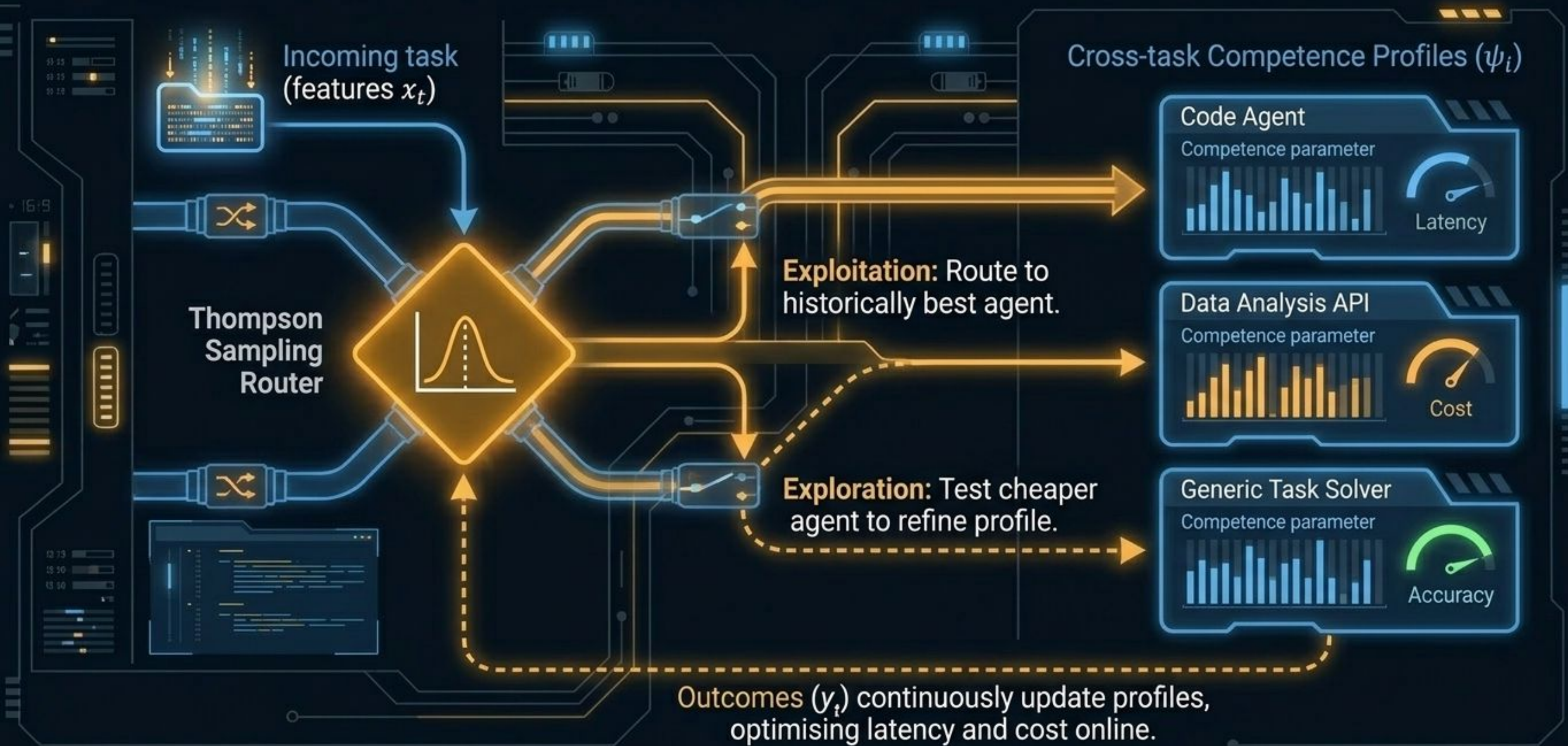
Case Study 1: Multi-Agent Code Generation and Testing



Case Study 2: Multi-Agent Incident Deliberation



Case Study 3: Bayesian Routing Across Task Ecosystems



A unified theory of operation across any agentic use case.



	 Code Generation	 Incident Deliberation	 Ecosystem Routing
The Latent Variable	Task Outcome ($Y \in \{0,1\}$)	Discrete Hypothesis ($h \in \mathcal{H}$)	Agent Competence (ψ_i)
The Evidence Source	Code Snippets & Test Logs	Conversational Arguments	Historical Success Labels
The Bayesian Action	Stop / Purchase more tests	Escalate to human / Conclude root cause	Select tool (Explore vs Exploit)

Regardless of the domain, the LLMs remain non-Bayesian generators.
The Bayesian structure dictates what to track, what to trust, and when to act.

Inside the Control State Core.

Task Outcomes & Hypotheses



Tracking the probability of ultimate success based on current evidence.

Cost & Risk Profiles

Dynamically updating budget allocation based on the asymmetric costs of errors.



Tool & Agent Reliability

Continuously modelling the trustworthiness of black-box LLMs and external APIs.



Error 0.0



Error 1.5



Error 3.0

**CONTROL
STATE
CORE**

Seven design patterns for Bayesian Orchestration.

1. Task-level belief factoring

Represents uncertainty in decision-relevant latent variables, not token uncertainty.

2. Messages as observations

Map agent text outputs to likelihood information via observation models.

3. Reliability-weighted updates

Temper likelihoods based on source bias to prevent overconfidence.

5. Utility-based control

Express routing and stopping as posterior expected-utility decisions.

4. Dependence-aware evidence pooling

Account for statistical dependence (e.g., shared prompts/retrieval).

7. Belief-state distillation

Compress interaction history into sufficient statistics to bypass context limits.

6. Cross-task posteriors

Adapt online routing by maintaining persistent tool profiles.

Known risks and engineering mitigations

 RISK	 ISSUE	 MITIGATION
  High-Dimensional Messages	Observation models mapping text to likelihoods can be misspecified	Continuous recalibration from measurable outcomes and likelihood tempering
  Correlated Evidence	Repeated tool calls to the same LLM yield dependent evidence, creating false confidence	Dependence-aware evidence pooling and conservative updating
  Computational Choke Points	Exact "Value-of-Information" calculations require expensive multi-step planning	Implement one-step approximations or learned surrogates for expected loss reduction

Takeaway: Bayesian control is resilient when updates are conservative and tied to verifiable outcomes

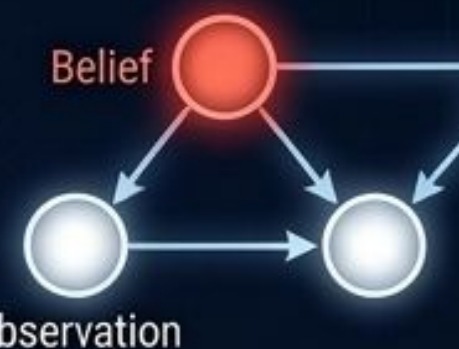
01. Benchmarking



Move beyond sequence prediction.
Implement outcome-based evaluation tracking predictive calibration and Value-of-Information (Vol) metrics.

Call to Action

02. Control layer modelling



Formalise belief states.
Learn observation models from outcome-labelled logs and recalibrate online.

03. Engineering & deployment



Confidence
Thresholds



Cost
Scales



Expose simple user controls (confidence/cost dials).
Distil interaction history to manage context windows.

04. Theory for orchestration



$$G(\mathbf{x}) = \prod_i \frac{P(x_i|\mathbf{x})}{\partial x_n}$$

$$G(\mathbf{e}) = \prod_i \frac{P(y_i|\mathbf{x})}{\partial x_n}$$

Develop decision-focused generalisation models.
Use Bayesian bandits for routing under non-stationarity.



The question is not whether an LLM 'is Bayesian'.
The question is whether your AI system behaves
as a **coherent decision maker** under uncertainty.

Let LLMs prioritise predictive capability. Let the orchestration layer prioritise utility, cost, and truth.