



[Home](#) [Licensing](#) [Patents](#) [Articles](#)

Inference Deployment Embodiments

by [Nick Clark](#) | Published March 27, 2026 | [PDF](#)

Three structural deployment configurations including embedded, co-resident, and hardware-assisted, providing identical governance guarantees with different latency and isolation profiles.

What It Is

Three structural deployment configurations including embedded, co-resident, and hardware-assisted, providing identical governance guarantees with different latency and isolation profiles.. This mechanism is defined in Chapter 8 of the cognition patent as a structural component of the agent's cognitive architecture, operating through deterministic evaluation rather than heuristic approximation.

Every aspect of this mechanism is specified declaratively in the agent's policy reference, making it auditable, reproducible, and governable without requiring access to the agent's internal decision-making process.

Why It Matters

Without inference deployment embodiments, inference proceeds without per-step governance. Current systems apply filtering only after generation is complete, missing the opportunity to prevent problematic inference trajectories before they produce harmful outputs. The fundamental distinction is between post-generation filtering and within-loop governance.

Post-generation filtering cannot prevent the generation of problematic content; it can only suppress it after resources have been consumed and semantic state has been contaminated by the generation process. Within-loop governance prevents problematic trajectories from developing in the first place, operating at a structurally different point in the inference pipeline.

How It Works Structurally

As defined in Chapter 8 of the cognition patent, inference deployment embodiments operates through a deterministic evaluation function embedded within the agent's cognitive architecture. The function receives structured inputs from the agent's canonical fields and produces outputs that govern subsequent processing stages. Every input, computation step, and output is recorded in the agent's lineage, ensuring complete reproducibility.

The semantic admissibility gate operates at each inference transition point. Before any candidate transition is committed, the gate evaluates it against the current semantic state, applicable policies, trust slope trajectory, and integrity constraints. The evaluation produces a deterministic admit, reject, or decompose decision. Rejected transitions are recorded as rejection events without affecting the semantic state.

What It Enables

This mechanism enables governed inference where every step is evaluated before commitment. Systems gain the ability to prevent problematic inference trajectories at the point of generation rather than filtering outputs after they have already been produced.

Because this mechanism is policy-governed and deterministic, it can be formally analyzed, audited, and certified. Regulatory compliance is demonstrable through structural analysis rather than solely through empirical testing. Different domains can tune the mechanism's parameters through policy configuration without requiring architectural changes, making the same structural capability applicable to autonomous vehicles, companion AI, therapeutic agents, and enterprise systems.

[Inference Control All 21 steps →](#)

Govern inference at the point of generation.

Primary Technical Disclosure

[◦ Inference-Time Semantic Execution Control](#)

Secondary Technical

[◦ Inference as Semantic Execution](#)◦ [Semantic Admissibility Gate](#)◦ [Entropy-Bounded Semantic Admissibility](#)◦ [Inference-Time Semantic Budget](#)◦ [Semantic Rollback and Checkpoint Recovery](#)◦ [Multi-Model Arbitration With Shared Semantic State](#)◦ [Structural Elegance Evaluation](#)◦ [Rights-Grade Inference Governance](#)◦ [Semantic State Object](#)◦ [Semantic State Object Schema](#)◦ [Inference Transition as Mutation](#)◦ [Trust-Slope Continuity Across Inference](#)◦ [Anchored Semantic Resolution](#)◦ [Semantic Lineage Recording](#)◦ [Policy-Governed Inference Execution](#)◦ [Partial State Handling](#)◦ [Model-Agnostic Inference Governance](#)◦ [Pre-Generation vs Post-Generation Distinction](#)◦ [Affect-Modulated Inference Admissibility](#)◦ [Integrity-Aware Inference](#)◦ [Confidence-Gated Inference Advancement](#)● [Inference Deployment Embodiments](#)

Applications (General)

[◦ Safety Without Alignment Theater: Why Structure Beats Supervision](#)◦ [How Commercial AI Platforms Reduce Prompt Size, Drift, and Governance Risk at Scale](#)◦ [When Execution Governance Becomes a Competitive Advantage — The Layer After LLM Gateways](#)◦ [Enterprise LLM Governance at the Point of Generation](#)◦ [Healthcare AI Admissibility Before Clinical Output](#)◦ [Inference Control for Legal Document Generation](#)◦ [Inference Control for Financial Advisory Output](#)◦ [Inference Control for Education Content Generation](#)◦ [Inference Control for Government Communications](#)

Applications (Specific)

[◦ Einstein Generates Without Semantic Admissibility](#)◦ [Databricks Serves Inference Without Semantic Gates](#)◦ [Snowflake Cortex Generates Without Admissibility Gates](#)◦ [Hugging Face Serves Models Without Semantic Governance](#)◦ [Cohere's Enterprise LLM Has No Semantic Admissibility Gate](#)◦ [Together AI Optimizes Inference Speed, Not Inference Governance](#)◦ [SageMaker Serves Models Without Semantic Admissibility](#)◦ [Vertex AI Generates Without Per-Transition Admissibility](#)◦ [Azure ML Deploys Models Without Admissibility Gates](#)◦ [Modal Runs Inference Fast Without Governing Output](#)◦ [Replicate Serves Open Models Without Semantic Governance](#)◦ [Fireworks AI Optimizes Speed Without Governing Semantics](#)◦ [Groq's LPU Accelerates Inference Without Governing It](#)◦ [Cerebras Achieves Wafer-Scale Inference Without Semantic Governance](#)

[Inference Control overview →](#)

AQ

deterministic

autonomy

Legal

Subject to one or more pending U.S. and international patent applications, see [Patents](#) for the current list and status. No license, express or implied, is granted. Any use requires a separate written agreement—see [Licensing](#). Patent applications referenced on this site are pending. Claim scope, if any, is subject to examination and may issue in altered form or not at all. See [Legal](#) for terms and conditions.

Adaptive Query™ is a trademark of Nicholas Clark. U.S. federal registration is pending. federal registration. AQ™, AQ Inside™, Adaptive Index™, Adaptive Network™, Semantic Agent™, @AQ™, AQID™, and Adaptive Coin™ are used as trademarks in connection with the Adaptive Query platform and brand. Other names may be trademarks of their respective owners.

Platform operated by Adaptive Query LLC, which provides patent and trademark licensing services. Copyright © 2025-2026 Nicholas Clark. All rights reserved.

Last updated: 2026-03-03



-
- [Inventive Steps](#)
- [Licensing](#)
- [Patents](#)
- [Articles](#)
- [Legal](#)
- [Opportunities](#)
- [Sitemap](#)



-
- nick@qu3ry.net
- 72 28 14 36 01



[Invented by Nick Clark](#) | Founding Investors: Devin Wilkie