

Coherence and Admissibility

FORMAL FOUNDATIONS FOR
GOVERNED ARTIFICIAL INTELLIGENCE



STRUCTURAL INTEGRITY | SEMANTIC COHERENCE | EXECUTABLE CONTROL

A formal framework for validating meaning, enforcing
constraints, and ensuring admissible AI system outputs.

Coherence and Admissibility

Formal Foundations for Governed Artificial Intelligence

Author: Michal Harcej

Date: 1st May 2026

Abstract

Artificial intelligence systems increasingly participate in workflows that produce legal, financial, medical, and safety-critical consequences. Existing evaluation methods focus primarily on correctness, confidence, safety, or alignment of model outputs. While valuable, these approaches do not determine whether an output is structurally permissible within a governed system.

This paper introduces a formal framework for **Coherence and Admissibility**, defining the conditions under which artificial intelligence outputs may participate in deterministic decision-making. We distinguish semantic coherence from governance admissibility, demonstrating that an output may be logically correct yet constitutionally inadmissible. We define coherence as consistency between observations, system state, and canonical knowledge, and admissibility as the satisfaction of executable governance constraints over proposed state transitions.

The framework establishes mathematical definitions for coherence spaces, admissibility predicates, constitutional invariants, governance rules, and deterministic execution semantics. We show that admissibility is a stronger property than correctness and that governed intelligent systems require admissibility verification prior to execution. The resulting architecture separates probabilistic reasoning from deterministic authority, enabling explainable, reproducible, and auditable intelligent systems.

1. Introduction

Artificial intelligence has become increasingly capable of producing outputs that rival or exceed human performance in numerous domains. However, correctness alone does not guarantee that an output should be executed.

A recommendation may be factually accurate while violating organizational policy. A prediction may exhibit high statistical confidence while conflicting with regulatory requirements. A planning system may generate an optimal strategy that exceeds authorized operational boundaries.

Existing evaluation frameworks largely assess model-centric properties such as accuracy, calibration, robustness, fairness, or alignment. These properties characterize the quality of generated outputs but do not determine whether those outputs are legally, operationally, or constitutionally executable.

This paper introduces a formal distinction between **coherence** and **admissibility**.

Coherence measures whether information is internally and externally consistent with system knowledge.

Admissibility determines whether coherent information may participate in authorized state transitions.

This distinction transforms AI governance from an interpretive discipline into a formally verifiable execution model.

2. Motivation

Modern intelligent systems implicitly assume the following sequence:

Input → AI → Decision → Execution

This architecture grants probabilistic components practical authority whenever generated outputs are executed without independent governance validation.

Governed systems require a different execution model:

Input → Observation → Coherence Verification → Admissibility Evaluation → Deterministic Decision → Execution

Under this model, execution depends not upon model confidence but upon formal admissibility.

3. Fundamental Definitions

Definition 1 (Observation)

An observation is a structured representation of measured system state.

Observations originate exclusively from authorized measurement components and possess no execution authority.

Definition 2 (Semantic Coherence)

Let

O denote the observation set,

K denote canonical knowledge,

S denote current system state.

Semantic coherence is defined as

$$C(O,K,S) \in [0,1]$$

where

$$C = 1$$

indicates complete consistency between observations, system state, and canonical knowledge.

$$C = 0$$

indicates complete inconsistency.

Definition 3 (Governance Rule)

A governance rule is an executable predicate

$$R_i : S \times A \rightarrow \{\text{True}, \text{False}\}$$

where

S denotes system state

A denotes proposed action.

Rules are deterministic.

Definition 4 (Admissibility)

An action a is admissible if and only if

$$\forall R_i \in G$$

$$R_i(s,a) = \text{True}$$

where G represents the complete governance rule set.

Admissibility therefore represents universal satisfaction of governing constraints.

Definition 5 (Constitutional Invariant)

An invariant

I

is a property that must remain true for every reachable system state.

Invariant violations immediately invalidate admissibility.

4. The Coherence Space

The system defines a coherence space

Ω

consisting of all observations compatible with canonical knowledge.

Formally,

$$\Omega = \{O \mid C(O,K,S) \geq \tau\}$$

where τ represents the minimum acceptable coherence threshold.

Observations outside Ω are considered constitutionally inconsistent.

No admissibility evaluation proceeds for incoherent observations.

5. Admissibility Predicate

Given

state s,

observation O,

candidate action a,

the admissibility predicate is

$$\text{Adm}(s,O,a)$$

defined by

$$\text{Adm} =$$

$$\text{Coherent}(O)$$

$$\wedge$$

$$\text{InvariantPreserved}(s,a)$$

$\wedge$

AuthorizedTransition(s,a)

$\wedge$

GovernanceSatisfied(s,a)

Only if

Adm=True

may execution proceed.

6. Execution Semantics

Execution follows six deterministic stages.

1. Observation acquisition.
2. Structural validation.
3. Coherence evaluation.
4. Governance resolution.
5. Admissibility evaluation.
6. Deterministic authorization.

Each stage is monotonic.

Failure terminates evaluation immediately.

No later stage may override an earlier rejection.

7. Admissibility Lattice

Outputs occupy one of four governance classes.

Admissible

Coherent and fully authorized.

Reviewable

Coherent but requiring external authorization.

Inadmissible

Violates one or more governance predicates.

Undefined

Cannot be evaluated because required information is absent.

Only the first category permits deterministic execution.

8. Properties

Property 1 (Correctness Does Not Imply Admissibility)

A correct output may violate governance constraints.

Therefore

Correctness $\nRightarrow$ Admissibility

Property 2 (Admissibility Implies Rule Satisfaction)

Every admissible action satisfies all governance predicates.

Property 3 (Incoherence Prevents Execution)

If

$C < \tau$

then

$Adm = \text{False}$.

Execution cannot proceed.

Property 4 (Deterministic Reproducibility)

Given identical observations,

identical governance rules,

identical system state,

the admissibility decision is identical.

9. Decision Trace

Every admissibility evaluation produces a structured trace containing:

- observation identifiers,
- coherence score,

- evaluated governance predicates,
- constitutional invariants,
- admissibility outcome,
- resulting state,
- execution authority.

Decision traces become proofs of admissibility rather than post hoc explanations.

10. Relationship to Existing Methods

Existing AI evaluation techniques emphasize prediction quality.

Formal verification demonstrates algorithmic correctness.

Policy engines enforce organizational rules.

The proposed framework integrates these perspectives by introducing admissibility as the final executable criterion governing state transitions.

Rather than replacing machine learning or formal verification, admissibility provides the constitutional interface between probabilistic reasoning and deterministic execution.

11. Applications

Potential application domains include:

- autonomous transportation,
- industrial control systems,
- financial transaction authorization,
- healthcare decision support,
- critical infrastructure,
- defense command systems,
- distributed protocol governance,
- safety-certified robotics.

Each domain benefits from separating predictive intelligence from execution authority.

12. Future Research

Future work includes:

- temporal coherence verification,
- distributed admissibility across federated systems,
- probabilistic coherence estimation,
- graph-theoretic governance resolution,
- admissibility-preserving protocol synthesis,
- formal proofs of invariant preservation,
- runtime verification over canonical knowledge graphs.

13. Conclusion

This paper introduced Coherence and Admissibility as complementary formal concepts for governed artificial intelligence.

Coherence establishes that observations are consistent with canonical knowledge and current system state. Admissibility establishes that coherent information may legally and architecturally participate in deterministic execution.

By distinguishing correctness from admissibility, the framework relocates governance from probabilistic model behavior to formal execution semantics. The resulting architecture enables reproducible authorization, explicit refusal, auditable decision traces, and mathematically verifiable governance, providing a foundation for intelligent systems whose legitimacy derives from enforceable structure rather than statistical confidence.