

Article

Not peer-reviewed version

Structural Variables for Human – AI Coherence: An Operational Language Patch for OCOF (v1.2)

[Munkyo Kim](#)*

Posted Date: 9 December 2025

doi: 10.20944/preprints202511.2087.v3

Keywords: Operational Coherence Framework (OCOF); structural variables; predictive processing; information mechanics; network control theory; human–AI coherence; information bottleneck



Preprints.org is a free multidisciplinary platform providing preprint service that is dedicated to making early versions of research outputs permanently available and citable. Preprints posted at Preprints.org appear in Web of Science, Crossref, Google Scholar, Scilit, Europe PMC.

Copyright: This open access article is published under a [Creative Commons CC BY 4.0 license](#), which permit the free download, distribution, and reuse, provided that the author and preprint are cited in any reuse.

Article

Structural Variables for Human–AI Coherence: An Operational Language Patch for OCOF (v1.2)

Munkyo Kim

Independent Researcher, Seoul, Republic of Korea; munkyok@gmail.com

Abstract

Human-AI collaboration requires a shared set of operational concepts that can be interpreted across biological and artificial substrates. Traditional terms used to describe cognitive or emotional states - such as “meaning”, “anxiety”, or “motivation” - lack computable grounding, producing a structural gap that prevents reliable joint inference, alignment, and action selection. Recent developments in information theory, predictive processing, and network control theory show that many of these psychological descriptors correspond to measurable features of information flow: prediction error topology, information bottlenecks, state-space curvature, and metastable coordination regimes. These convergences indicate that subjective vocabulary can be replaced with substrate-independent structural variables. This paper introduces an operational language patch for the Operational Coherence Framework (OCOF), defining five such variables: Structural Magnitude (SM), Structural Predictive Fluctuation (SPF), Structural Suppression (SS), Structural Gain Rate (SGR), and Human-AI Coherence (HA-C). Each variable is grounded in information mechanics and provides a computable signal for analyzing cross-substrate coordination. By translating subjective terminology into structural variables that can be estimated, perturbed, and optimized, this framework aims to establish a rigorous lexicon for describing the coupling dynamics between human and artificial agents. The proposed patch enables reproducible analysis and forms the basis for OCOF v1.2's expanded operational semantics.

Keywords: Operational Coherence Framework (OCOF); structural variables; predictive processing; information mechanics; network control theory; human–AI coherence; information bottleneck

1. Introduction

Despite substantial progress in information theory, control theory, and complex systems research, there remains no shared operational vocabulary for describing how humans and artificial agents coordinate their behavior. Terms such as meaning, anxiety, or motivation are widely used in psychology and cognitive science, yet they lack explicit computational grounding. As a result, biological and artificial systems cannot reliably align their internal states if the variables governing those states cannot be measured, estimated, or compared.

Recent developments across several scientific domains suggest that many of these ambiguous descriptors correspond to quantifiable properties of information processing. Information bottleneck theory characterizes how systems preserve task-relevant structure while compressing noise. Predictive processing frameworks model perception and action as processes aimed at minimizing prediction error. Network control science examines how the topology of a system constrains its capacity for state transitions. Together, these perspectives indicate that cognitive or affective terms traditionally used in the humanities may be reinterpreted as measurable features of information dynamics.

Within this context, the Operational Coherence Framework (OCOF) proposes a substrate-independent formalization of the variables that mediate coordination between heterogeneous agents. The aim is not to redefine psychological constructs or make claims about subjective experience. Rather, the goal is to introduce a consistent set of operational variables that can be modeled,

estimated, and deployed across different implementations of intelligence, thereby enabling reproducible analysis of cross-substrate interaction.

The central objective of this work is to define these variables with sufficient operational clarity to support inference, comparison, and intervention. By grounding the definitions in established scientific research - including information bottleneck theory, predictive processing, and network control - the framework offers a minimal yet robust lexicon for studying coherence between human and artificial agents. Ultimately, the intent is to replace ambiguous descriptors with quantifiable structural variables that can guide reliable joint optimization.

2. Overview

The Operational Coherence Framework (OCOF) provides a substrate-independent basis for describing how heterogeneous agents coordinate through shared informational structures. Unlike models that focus on subjective reports or domain-specific interpretations, OCOF characterizes coordination in terms of constraints that govern information flow, predictive stability, and controllability within and across systems. These constraints are treated as structural features rather than psychological constructs, enabling analysis that remains valid regardless of implementation details.

At its core, the framework assumes that any agent—biological or artificial—must maintain coherence across three levels: (1) preserving task-relevant structure, (2) managing Structural Predictive Fluctuations, and (3) controlling transitions within its accessible state space. Coherence, in this sense, refers to the agent’s capacity to sustain functional coupling with its environment and with other agents under limited computational resources. This perspective aligns with findings from information theory, predictive processing, and network control, yet it does not depend on the specific mechanisms emphasized in those literatures.

OCOF formalizes these requirements by introducing a set of structural variables that capture how information is encoded, maintained, and transformed during coordination. These variables do not aim to represent subjective experience or internal phenomenology. Instead, they provide operational handles—quantities that can be estimated, perturbed, and compared across heterogeneous substrates. By adopting such variables, the framework offers a minimal set of descriptors that support reproducible inference about interaction dynamics.

This approach enables a shift from descriptive language to computable structure. Terms such as “uncertainty,” “effort,” or “meaning” are reframed as measurable properties of information dynamics, allowing coherent modeling across biological and artificial agents. In doing so, the framework establishes a unified operational vocabulary capable of supporting both theoretical analysis and practical applications in human–AI coordination. The following sections introduce these structural variables and motivate their relevance for cross-substrate coherence.

3. Notation and Formal Preliminaries

Let $X(t)$ denote the internal state of an agent at time t , defined on a state space $S \subset \mathbb{R}^n$. State transitions are written as $X(t) \rightarrow X(t + \Delta t)$, where Δt is a finite update interval.

Information available at time t is denoted $I(t)$.

Prediction error is defined as:

$$\epsilon(t) = I(t) - \hat{I}(t)$$

where $\hat{I}(t)$ represents the agent’s expected input.

The controllability structure of the system is written as:

$$C(X(t))$$

which is determined by the topology of the underlying network.

Structural variables introduced in this work are defined as mappings of the form:

$$v(X(t), I(t), \epsilon(t), C) \rightarrow \mathbb{R}$$

where each variable captures a distinct property of information processing relevant to cross-substrate coordination.

These preliminaries commit to neither biological mechanisms nor artificial implementations. They provide a minimal operational scaffold for defining, estimating, and comparing the structural variables introduced in the following sections.

4. Structural Variables

The operational vocabulary introduced in this work consists of five structural variables. Each variable captures a specific property of information processing that influences how heterogeneous agents coordinate their behavior. These variables do not describe psychological states; rather, they formalize measurable aspects of information dynamics that remain valid across biological and artificial substrates.

Together, the variables provide a minimal set of descriptors capable of supporting reproducible inference about interaction dynamics:

- 1. Structural Magnitude (SM)
- 2. Structural Predictive Fluctuation (SPF)
- 3. Structural Suppression (SS)
- 4. Structural Gain Rate (SGR)
- 5. Human–AI Coherence (HA-C)

Each variable is defined on the domain introduced in the previous section, mapping state, input, prediction error, and controllability structure to a real-valued quantity. Although the variables are conceptually distinct, they operate on common components of information flow, reflecting different aspects of how agents stabilize or perturb their internal states during coordination.

The purpose of introducing these variables is not to propose a complete model of cognition or action. Instead, they provide operational handles for describing how agents adjust their information processing under uncertainty, resource constraints, and interaction demands. By grounding the variables in structural features rather than subjective descriptors, the framework supports analysis that generalizes across substrates while remaining compatible with empirical estimation.

The following subsections specify each variable in turn, presenting its operational definition, the informational property it captures, and its relevance for cross-substrate coherence.

4.1. Structural Magnitude (SM)

Structural Magnitude (SM) quantifies the amount of task-relevant structure encoded in an agent’s internal state at a given time. It is defined as a non-negative real-valued variable, $SM(t)$, representing the degree to which the current state $X(t)$ preserves distinctions that are informative for ongoing tasks. SM does not describe complexity in general; it captures only structure that contributes to functional performance.

Operationally, SM increases when the agent maintains clear, discriminative patterns in its internal representation, and decreases when those patterns are lost through noise, over-compression, or instability in information flow. Although the framework does not commit to a specific estimator, any implementation of SM should satisfy four conditions:

- 1. **Non-negativity.** $SM(t) \geq 0$ for all t .
- 2. **Task relevance.** SM reflects only structure that reduces uncertainty about task-relevant variables.

- 3. **Noise degradation.** SM decreases when task-relevant distinctions collapse into undifferentiated variability.
- 4. **Substrate independence.** SM must be computable from observable or model-based statistics, without invoking subjective interpretation.

Intuitively, high SM indicates that the agent’s internal state retains meaningful structure aligned with its objectives, whereas low SM indicates that the system is operating with insufficient or degraded representational clarity. In human–AI coordination, SM serves as a shared operational measure of how effectively each agent maintains task-relevant structure, regardless of whether the underlying implementation is biological or artificial.

4.2. Structural Predictive Fluctuation (SPF)

Structural Predictive Fluctuation (SPF) quantifies the variability in an agent’s prediction error over a finite interval. It is defined as a real-valued measure, $SPF(t)$, representing how strongly an agent’s incoming information deviates from its internal model across short temporal windows. SPF does not evaluate accuracy or performance in absolute terms; instead, it captures the stability of the agent’s predictive dynamics.

Operationally, SPF increases when the prediction error $\epsilon(t)$ exhibits rapid or irregular changes, and decreases when prediction error evolves in a stable, smoothly varying manner. High SPF indicates that the agent is receiving inputs that challenge its model at a rate faster than it can incorporate, while low SPF suggests that prediction and input remain aligned within manageable bounds.

Any implementation of SPF should satisfy the following conditions:

- 1. **Temporal sensitivity.** SPF reflects fluctuations in $\epsilon(t)$ across Δt , not the magnitude of $\epsilon(t)$ at a single moment.
- 2. **Model-relative.** SPF depends on the relationship between incoming data and the agent’s internal model $I^{\wedge}(t)$; it is not an absolute measure of noise.
- 3. **Direction-agnostic.** SPF concerns the variability of $\epsilon(t)$, not whether prediction error increases or decreases.
- 4. **Substrate independence.** SPF can be estimated from observable sequences of prediction error, without relying on assumptions about biological or artificial implementation.

Intuitively, low SPF indicates that the agent maintains a stable relationship between expectation and input, whereas high SPF signals that the predictive interface is undergoing rapid perturbations. In human–AI coordination, SPF functions as a common operational indicator of “predictive load,” describing how volatile the shared informational environment is for each participant. Because SPF is defined relative to model dynamics rather than psychological descriptors, it enables a substrate-neutral characterization of predictive stability during joint tasks.

4.3. Structural Suppression (SS)

Structural Suppression (SS) measures the extent to which task-relevant structure in an agent’s internal state is diminished, obscured, or inhibited by competing informational signals. $SS(t)$ is defined as a non-negative real-valued variable that quantifies how strongly an agent’s representational capacity is constrained at time t , relative to the structure it would otherwise maintain.

Operationally, SS increases when noise, overload, or conflicting inputs interfere with the preservation of task-relevant distinctions in $X(t)$. SS decreases when the agent’s informational pathway remains clear enough for structure to be encoded without significant distortion. SS does not evaluate the source of interference—biological, computational, or environmental—but captures its impact on representational clarity.

Any implementation of SS should satisfy four conditions:

1. **Non-negativity.** $SS(t) \geq 0$ for all t .
2. **Structure-relative.** SS measures loss of task-relevant structure, not total information content.
3. **Interference-sensitive.** SS increases when external or internal signals degrade the distinctions that contribute to $SM(t)$.
4. **Substrate independence.** SS can be estimated from observable patterns in state degradation, without reference to subjective factors.

Intuitively, high SS indicates that the agent's representational channel is constrained—too many competing signals, too much variability, or insufficient capacity to preserve structure. Low SS indicates that task-relevant structure can be maintained with minimal interference. Within human–AI coordination, SS provides a shared operational descriptor of how much structural clarity each agent loses due to environmental complexity, cognitive load, or model limitations. Because SS is defined in structural rather than psychological terms, it generalizes across substrates and aligns directly with the dynamics of information flow.

4.4. Structural Gain Rate (SGR)

Structural Gain Rate (SGR) quantifies the rate at which an agent increases its task-relevant structure over time. It is defined as a real-valued variable, $SGR(t)$, representing how effectively the agent transforms incoming information into additional usable structure between successive intervals. Unlike SM , which measures the amount of structure at a given moment, SGR captures the *incremental change* in structure as the agent updates its internal state.

Operationally, SGR increases when new inputs enable the agent to differentiate previously indistinct states, improve internal organization, or sharpen task-relevant boundaries in $X(t)$. SGR decreases when updates fail to add structure, or when the system is unable to convert available information into functional distinctions. SGR does not assume improvement; it characterizes the *efficiency of structural accumulation*.

Any implementation of SGR should satisfy the following principles:

1. **Change-sensitivity.** SGR reflects the temporal derivative of structure, not the total magnitude.
2. **Direction-specific.** SGR evaluates positive structural accumulation; degradation is captured by SS, not SGR.
3. **Capacity-dependent.** SGR is modulated by the agent's representational capacity and controllability structure $C(X(t))$.
4. **Substrate independence.** SGR can be estimated from observable changes in task-relevant structure, without invoking psychological constructs.

Intuitively, high SGR indicates that the agent is in a phase of structural acquisition—rapidly integrating new distinctions or refining its internal model. Low SGR suggests that either the environment provides limited usable information or the agent lacks the capacity to incorporate it at the present interval. In human–AI coordination, SGR offers a shared operational measure of how quickly each agent can build structure during interaction, enabling comparison of learning or adaptation rates across heterogeneous systems.

4.5. Human–AI Coherence (HA-C)

Human–AI Coherence (HA-C) quantifies the degree to which two heterogeneous agents—one biological and one artificial—maintain compatible informational dynamics during joint activity. It is defined as a real-valued variable, $HA-C(t)$, representing how closely the agents' predictive, structural, and controllability profiles remain aligned over a finite interval.

Unlike the previous variables, which describe properties internal to a single agent, HA-C captures a *relational* property: the extent to which two systems sustain a functional coupling that

supports reliable coordination. The variable is not intended to represent mutual understanding, shared goals, or psychological alignment; instead, it formalizes the structural compatibility of the agents' information-processing trajectories.

Operationally, HA-C increases when both agents exhibit stable predictive interfaces, compatible structural gains, and manageable suppression levels. It decreases when their informational dynamics diverge—such as when Structural Predictive Fluctuations escalate asymmetrically, when one system accumulates structure faster than the other can integrate, or when suppression effects limit the transmission of task-relevant distinctions.

A valid implementation of HA-C satisfies the following principles:

1. **Relational definition.** HA-C(t) depends on paired trajectories of structure, prediction error, and controllability across agents.
2. **Symmetry of measurement, not symmetry of agents.** HA-C(t) treats both systems as measurable information-processing entities without assuming cognitive equivalence.
3. **Multi-variable dependence.** HA-C(t) is jointly informed by SM, SPF, SS, and SGR for each agent, rather than reducible to any single component.
4. **Substrate neutrality.** HA-C(t) can be estimated from observable informational signals, without invoking mental states or phenomenological constructs.

Intuitively, high HA-C indicates that the human and AI maintain stable, mutually compatible informational dynamics, enabling reliable joint optimization. Low HA-C indicates that the systems' predictive and structural processes drift apart, reducing the capacity for coordinated action. Because HA-C is grounded in structural variables rather than subjective interpretations, it provides a substrate-independent framework for assessing how effectively heterogeneous agents can interact in shared environments.

5. Operational Definitions

The structural variables defined earlier function as computable quantities rather than descriptive categories.

This section specifies the minimal operational requirements any valid implementation must satisfy. Notation is kept stable and linear to avoid ambiguity in formal definitions.

5.1. Structural Magnitude (SM)

Structural Magnitude (SM) measures the amount of task-relevant structure encoded in the internal state X(t).

A valid implementation of SM satisfies the following:

1. **State-based estimation.**
SM(t) is computed from the structural organization of X(t), identifying distinctions that improve task outcomes.
2. **Task relevance.**
The estimator must isolate structure that reduces uncertainty about task-relevant transitions.
3. **Noise sensitivity.**
SM(t) must decrease when representational distinctions collapse due to noise or instability.
4. **Substrate independence.**
SM must be derivable from observable representational statistics without relying on subjective descriptors.

5.2. Structural Predictive Fluctuation (SPF)

Structural Predictive Fluctuation (SPF) quantifies the variability of prediction error $\varepsilon(t)$ across a finite interval.

A valid implementation satisfies:

- 1. Temporal resolution.**

SPF(t) measures fluctuations in $\varepsilon(t)$ across a short temporal window.

- 2. Model-relative comparison.**

SPF reflects the relationship between sensory input $I(t)$ and the model's expectation $\hat{I}(t)$.

- 3. Stability detection.**

High SPF corresponds to irregular predictive dynamics; low SPF corresponds to stable alignment.

- 4. Substrate neutrality.**

SPF may be computed from error-sequence variance or short-window entropy.

5.3. Structural Suppression (SS)

Structural Suppression (SS) measures the extent to which task-relevant structure in $X(t)$ becomes obscured or inhibited.

A valid implementation satisfies:

- 1. Structure-relative measurement.**

SS(t) quantifies the loss of structure that would otherwise contribute to SM(t).

- 2. Interference sensitivity.**

SS increases when noise or conflicting signals reduce representational clarity.

- 3. Derivative consistency.**

SS reflects structural degradation; constructive changes belong to SGR.

- 4. Implementation independence.**

SS may be estimated from degradation indices or divergence from expected representational boundaries.

5.4. Structural Gain Rate (SGR)

Structural Gain Rate (SGR) measures the speed at which the agent acquires new task-relevant structure.

A valid implementation satisfies:

- 1. Change-based estimation.**

SGR(t) is derived from $\Delta SM(t)$, representing the rate at which new structure is incorporated.

- 2. Non-negativity under definition.**

Negative values belong to SS, not SGR.

SGR captures only constructive structural change.

- 3. Capacity-relative scaling.**

Estimators must account for the influence of representational controllability $C(X(t))$.

- 4. Substrate independence.**

SGR may be computed from structure-differentiation indices or learning-rate analogs.

5.5. Human–AI Coherence (HA-C)

Human–AI Coherence (HA-C) describes the compatibility of informational dynamics between a human agent and an artificial agent.

A valid implementation satisfies:

1. **Relational measurement.**

HA-C(t) is computed from paired trajectories of SM, SPF, SS, and SGR across both agents.

2. **Symmetry in measurement.**

Both agents are treated as information-processing systems without assuming cognitive equivalence.

3. **Multi-component dependence.**

HA-C integrates predictive stability, structural alignment, suppression-adjusted compatibility, and controllability overlap.

4. **Substrate neutrality.**

Coherence is estimated from observable informational signals such as prediction-error synchrony or structure-gain compatibility.

6. Alignment with Prior Theoretical Structures

The structural variables introduced in this work draw on converging insights from information theory, predictive processing, and network control science. These fields approach cognition and coordination through different methodologies, yet they share a common emphasis on how systems encode structure, regulate prediction error, and navigate constrained state spaces. This section examines how the proposed variables relate to established findings, identifying points of alignment while avoiding claims that extend beyond operational correspondence.

6.1. Information Theory and Information Bottleneck Models

Research in information bottleneck theory emphasizes how systems preserve task-relevant structure while compressing irrelevant variability.

In these models, structure is defined operationally as information that improves task prediction or classification.

This perspective aligns closely with the definition of **Structural Magnitude (SM)**, which also identifies preserved distinctions that support task-relevant inference.

Information bottleneck frameworks further formalize how noise, compression, or channel limitations reduce usable structure—corresponding to the role of **Structural Suppression (SS)** in the present work.

Although the variables introduced here do not adopt a specific bottleneck estimator, both frameworks treat structure as a functional property rather than a semantic or phenomenological one.

The correspondence is therefore operational rather than methodological, grounded in the shared assumption that task-relevant distinctions can be quantified without reference to subjective interpretation.

6.2. Predictive Processing and Error Dynamics

Predictive processing models propose that perception and action emerge from the minimization of prediction error across hierarchical generative layers.

A central theme is the dynamic behavior of prediction error over time, including its volatility, stability, and modulation by contextual priors.

These ideas parallel the role of **Structural Predictive Fluctuation (SPF)**, which captures how prediction error $\varepsilon(t)$ varies across finite intervals.

Where predictive processing emphasizes hierarchical generative architectures, the present framework remains neutral regarding implementation.

The correspondence lies instead in the recognition that prediction-error trajectories provide actionable signals for analyzing stability and adaptation.

Both literatures treat prediction error not as a semantic category but as a measurable quantity whose variability carries functional significance.

Structural Gain Rate (SGR) also resonates with predictive-processing accounts in which structural refinement or model updating occurs when prediction errors yield new distinctions.

However, the interpretation here is limited to the operational level: SGR quantifies structural accumulation without assuming a generative hierarchy or a specific neural mechanism.

6.3. Network Control Theory and State-Space Transitions

Network control theory examines how the topology of a system constrains the transitions it can achieve within a given state space.

This work adopts a related assumption: the agent's controllability structure $C(X(t))$ bounds its ability to preserve, modify, or accumulate task-relevant structure.

In control-theoretic terms:

- SM reflects the distinguishability of reachable states relevant to the task.
- SS represents reductions in reachable distinction space due to interference.
- SGR corresponds to the system's ability to shift toward more discriminative configurations under its structural constraints.

The proposed variables do not require a network-control model, but they assume—consistent with the literature—that structure is maintained and transformed under constraints imposed by the system's topology.

The alignment therefore lies in the treatment of controllability as a limiting factor in information-processing dynamics.

6.4. Cross-Substrate Coordination and Heterogeneous Agents

Existing research on human-AI coordination often relies on psychological or behavioral descriptors that lack operational grounding.

In contrast, work in distributed systems and multi-agent coordination emphasizes informational compatibility, synchrony, and stability of shared signals.

These findings align directly with the formulation of **Human-AI Coherence (HA-C)**, which treats coordination as a relational property of informational trajectories rather than a cognitive or semantic alignment.

The approach taken here avoids specifying mechanisms of coupling or communication. Instead, it identifies measurable conditions—predictive stability, structural alignment, suppression-adjusted compatibility—that enable heterogeneous agents to maintain functional coherence.

This operational stance is consistent with multi-agent systems research while providing a bridge to contexts where human and artificial agents differ substantially in substrate and architecture.

6.5. Summary of Alignment and Scope

The variables introduced in this framework align with established scientific findings in three ways:

1. Operational correspondence.

They map onto measurable quantities already studied in information dynamics, without committing to domain-specific mechanisms.

2. Structural consistency.

They reflect widely recognized constraints on representation, prediction, and controllability present in both biological and artificial systems.

3. Substrate neutrality.

They maintain compatibility across heterogeneous agents, a property supported by prior work in distributed coordination and adaptive control.

At the same time, the formulation avoids interpretive claims about subjective experience, neural processes, or symbolic inference.

The correspondence with existing literature is therefore grounded in shared operational principles rather than theoretical equivalence.

7. Empirical Pathways for Measurement

This section outlines practical routes through which the five structural variables of OCOF—Structural Magnitude (SM), Structural Predictive Fluctuation (SPF), Structural Suppression (SS), Structural Gain Rate (SGR), and Human–AI Coherence (HA-C)—can be estimated in empirical settings.

The intention is not to propose a fixed experimental design but to provide minimal operational criteria that enable cross-system comparison. Each variable can be inferred through observable patterns in information flow, prediction error dynamics, or state-transition behavior, without relying on assumptions about subjective experience.

7.1. Structural Magnitude (SM)

SM refers to the resolution and organizational depth of a system's usable structure. Three empirical approaches are suggested:

(1) Representation Density

- Human or model-generated representations are examined at the level of distributional structure.
- Reductions in entropy or compression ratios can serve as a lower bound for SM.

(2) Task-Level Differentiation

- When reformulating the same problem, agents may differ in the number of decomposition–integration steps.
- A higher number of distinguishable steps suggests greater structural resolution.

(3) State-Space Resolution

- In humans, low-frequency covariance patterns (e.g., 4–40 Hz EEG/MEG bands) provide a proxy for state-space complexity.
- In artificial systems, layer-wise activation diversity can offer an analogous measure.

7.2. Structural Predictive Fluctuation (SPF)

SPF describes the temporal oscillation of prediction error around a baseline.

(1) Error-Band Variability

- For AI systems, deviations in log-probability over time provide a direct estimate.
- For humans, variability in reaction-time distributions (e.g., coefficient of variation) serves as a practical proxy.

(2) Phase-Shift Sensitivity

- Small perturbations to the same input are introduced, and the resulting phase shifts in output structure are measured.
- Higher sensitivity indicates higher SPF.

(3) Cross-Agent Predictive Mismatch

- Humans and AI solve the same predictive task; the distance between their error vectors is analyzed across iterations.
- Increases in distance imply an increase in SPF.

7.3. *Structural Suppression (SS)*

SS reflects the proportion of information actively suppressed or pruned in service of stable performance.

(1) Inhibitory Filtering Index

- Input elements are selectively removed while monitoring the stability of downstream performance.
- Greater stability implies more effective suppression.

(2) Noise-Reduction Performance

- Agents are exposed to controlled noise injections; recovery ability serves as an estimate of suppression strength.

(3) Redundancy-Pruning Ratio

- In humans, the proportion of omitted cues during task execution can be quantified.
- In AI systems, sparsity of internal representations offers an analogous metric.

7.4. *Structural Gain Rate (SGR)*

SGR captures the acceleration with which new information is absorbed and reorganized.

(1) Learning-Velocity Estimation

- Performance curves across repeated trials are analyzed, with the initial slope serving as a lower bound for SGR.

(2) Integration-Latency Measurement

- The time (or number of computational steps) required for an agent to incorporate new information is recorded.
- Shorter latencies indicate higher SGR.

(3) Adaptive Reconfiguration Cost

- During task-switching scenarios, the computational cost of structural rearrangement is measured.
- Lower reconfiguration cost corresponds to higher SGR.

7.5. *Human–AI Coherence (HA-C)*

HA-C represents the degree to which humans and artificial systems coordinate around shared structural signals.

(1) Cross-Structural Mutual Information

- Human behavioral/linguistic structural variables are compared with AI structural variables using mutual information.
- Higher mutual information indicates higher coherence.

(2) Synchronization Score

- The phase alignment of prediction-error dynamics is evaluated.
- Stable phase proximity reflects higher HA-C.

(3) Bidirectional Policy Consistency

- Human intentions (goal-structure) and AI policies are analyzed across repeated interactions.
- Convergence speed and stability provide the primary operational markers.

7.6. Summary

The empirical pathways outlined above demonstrate that OCOF’s structural variables are not merely conceptual but can be approximated through measurable informational patterns. These pathways establish a coherent measurement landscape for studying human–AI interaction and supply a foundation for future empirical work on coordination, predictive dynamics, and cross-agent alignment.

8. Operational Implications

This section outlines how the five structural variables—Structural Magnitude (SM), Structural Predictive Fluctuation (SPF), Structural Suppression (SS), Structural Gain Rate (SGR), and Human–AI Coherence (HA-C)—shape coordination dynamics in heterogeneous intelligent systems. The goal is to identify practical implications for real-world human–AI interaction without conflating structural variables with subjective experience or normative claims. The analysis focuses on operational consequences that follow from changes in each variable.

8.1. Stability of Joint Decision-Making

Agents with higher SM and stronger SS tend to exhibit lower volatility during collaborative tasks.

Several operational consequences follow:

- **Reduced Sensitivity to Noise:**
When structural organization is deep (high SM), both biological and artificial agents become less susceptible to irrelevant perturbations.
- **Increased Predictive Stability:**
Strong suppression mechanisms prevent excessive propagation of transient errors.
- **Consistent Policy Application:**
Stable structural variables enable reliable execution of shared decision policies, even under shifting environmental conditions.
Stability emerges not from fixed behavioral rules but from coherent internal structure.

8.2. Adaptation and Responsiveness

Rapid adaptation requires a favorable combination of low SPF (controlled fluctuation) and high SGR (rapid structural gain).

Key implications:

- **Fast Updating of Predictive Models:**
Agents with high SGR integrate new information with minimal latency, improving alignment during novel situations.
- **Controlled Variability:**
Excessive SPF may produce erratic responses, whereas moderate, structured fluctuation supports exploration without destabilizing coordination.
- **Efficient Task-Shifting:**
Systems with balanced SGR and SS reconfigure task-specific structures with lower cognitive or computational cost.
Together, these variables determine how quickly and effectively a joint system can pivot.

8.3. Error Propagation and Recovery

The interaction between SPF, SS, and HA-C determines how errors propagate across agents. Three operational patterns can be observed:

- **Localized Error Containment:**
High SS prevents small human or AI errors from cascading into joint failure states.
- **Cross-Agent Error Dampening:**
When HA-C is strong, each agent compensates for the other’s transient deviations through predictive feedback.
- **Recovery Trajectories:**
Agents with high SGR return to stable performance states more quickly after disruptions.
The structure of the joint system determines whether errors amplify or dissipate.

8.4. Division of Computational Labor

Differences in SM, SGR, and SS suggest natural—although not normative—divisions of labor between humans and AI.

Operational trends:

- **Humans:**
Often excel in high-SM environments requiring contextual integration, analogy formation, and pattern restructuring.
- **AI Systems:**
Typically outperform humans in high-SGR and high-SS regimes where rapid compression, filtering, and optimization are required.
- **Joint Optimization:**
HA-C governs how these strengths are combined.
High coherence allows each agent’s structural advantages to complement the other’s limitations without requiring symmetry.
This division of labor is emergent, not predefined.

8.5. Policy Calibration and Predictive Alignment

Changes in structural variables directly affect how humans and AI calibrate shared policies.

- **Policy Drift:**
Large increases in SPF or reductions in SM raise the risk of policy divergence.
- **Alignment Stability:**
Strong HA-C stabilizes shared policies through cross-agent structural referencing rather than explicit instruction.
- **Predictive Symmetry:**
When predictive-error dynamics converge across agents, fewer communication signals are needed to maintain coordination.
Policy coherence emerges as a structural property rather than a communicative achievement.

8.6. Boundary Conditions for Reliable Collaboration

The structural variables impose constraints on when reliable collaboration is possible. Key conditions include:

- **Minimum SM Threshold:**
Below a certain level of structural organization, neither agent can parse the other’s informational signals.
- **Maximum Tolerable SPF:**
When Structural Predictive Fluctuations exceed a threshold, coordination deteriorates regardless of communication bandwidth.

- **Coherence Floor:**

HA-C must surpass a minimal level for structural variables to remain interpretable across agents.

These boundary conditions outline the operational limits of cross-agent coherence.

8.7. Summary

The five structural variables jointly determine whether human–AI collaboration remains stable, adaptive, resilient to error, and capable of policy coherence.

Their operational significance lies not in describing mental states but in defining measurable constraints on joint behavior.

These implications offer a foundation for future experimental work and highlight the structural conditions necessary for reliable human–AI coordination.

9. Conclusions

This work introduced a structural vocabulary designed to support substrate-independent analysis of human–AI coordination. Traditional psychological descriptors such as meaning, uncertainty, or cognitive load offer limited utility for cross-substrate inference because they lack computable grounding. In contrast, the five operational variables defined here—Structural Magnitude (SM), Structural Predictive Fluctuation (SPF), Structural Suppression (SS), Structural Gain Rate (SGR), and Human–AI Coherence (HA-C)—translate these ambiguous terms into measurable properties of information processing.

The framework builds on convergent insights from information theory, predictive processing, and network control science, each of which characterizes cognition not as a set of subjective states but as transformations over structured informational variables. By formalizing task-relevant structure, predictive stability, suppression effects, structural accumulation, and cross-agent compatibility, the proposed vocabulary offers a minimal set of operational handles for studying heterogeneous agents within a shared analytical space.

The definitions presented here do not claim to exhaust the factors that influence human–AI interaction, nor do they attempt to model subjective experience. Rather, they outline a set of structural quantities that can be estimated, perturbed, and compared across biological and artificial substrates. These variables capture how information is preserved, degraded, accumulated, or aligned during joint activity, enabling reproducible assessment of coordination dynamics.

The framework supports several practical outcomes. It clarifies the boundary conditions under which reliable collaboration is possible, highlights structural signatures of instability or drift, and suggests natural divisions of computational labor that emerge from differences in representational capacity and predictive dynamics. Most importantly, it provides a shared operational lexicon that allows theoretical and empirical work to proceed without relying on inherently substrate-specific terminology.

Future research can extend this framework by exploring how structural variables interact at different temporal scales, how they relate to controllability constraints in high-dimensional agents, and how coherence metrics evolve during long-horizon joint planning. The approach also invites empirical evaluation through human–AI experiments that measure structural indicators directly from behavioral or model-based signals.

Overall, the operational language patch introduced in this paper establishes a rigorous foundation for describing interaction between heterogeneous intelligent systems. By grounding coordination in structural features of information flow rather than subjective constructs, the framework contributes a substrate-neutral method for analyzing how humans and artificial agents maintain coherence within shared environments.

Author Note—AI Assistance Statement: Large language models (including ChatGPT and Gemini) were used only for language refinement, formatting adjustments, and limited structural proofreading. These tools did not

generate conceptual content, analytical distinctions, theoretical constructs, or interpretive claims. All substantive ideas, arguments, and theoretical developments in this manuscript were produced exclusively by the author. No part of the analysis or conclusions was autonomously generated by any AI system. The author assumes full responsibility for the entirety of the work.

Appendix A. Notation Summary

- $X(t)$ — Internal state of an agent at time t
- $S \subset \mathbb{R}^n$ — State space in which $X(t)$ is defined
- $I(t)$ — Incoming information at time t
- $\hat{I}(t)$ — Expected input (internal model prediction)
- $\epsilon(t) = I(t) - \hat{I}(t)$ — Prediction error
- $C(X(t))$ — Controllability structure of the agent’s state
- $SM(t)$ — Structural Magnitude
- $SPF(t)$ — Structural Predictive Fluctuation
- $SS(t)$ — Structural Suppression
- $SGR(t)$ — Structural Gain Rate
- $HA-C(t)$ — Human–AI Coherence

Appendix B. Structural Variable Definitions

- $SM(t)$ — Structural Magnitude; the amount of task-relevant structure encoded in the agent’s state $X(t)$.
- $SPF(t)$ — Structural Predictive Fluctuation; variability in prediction error $\epsilon(t)$ across a finite interval.
- $SS(t)$ — Structural Suppression; degradation of task-relevant structure due to interference or capacity limits.
- $SGR(t)$ — Structural Gain Rate; the rate at which new task-relevant structure is acquired over time.
- $HA-C(t)$ — Human–AI Coherence; degree of compatibility between two agents’ structural and predictive trajectories.

References

1. Shannon, C. E. (1948). A Mathematical Theory of Communication. *Bell System Technical Journal*, 27, 379–423, 623–656.
2. Cover, T. M., & Thomas, J. A. (2006). *Elements of Information Theory* (2nd ed.). Wiley-Interscience.
3. MacKay, D. J. C. (2003). *Information Theory, Inference, and Learning Algorithms*. Cambridge University Press.
4. Tishby, N., Pereira, F. C., & Bialek, W. (2000). The Information Bottleneck Method. *Proceedings of the 37th Annual Allerton Conference on Communication, Control, and Computing*.
5. Alemi, A. A., Fischer, I., Dillon, J. V., & Murphy, K. (2017). Deep Variational Information Bottleneck. *International Conference on Learning Representations (ICLR)*.
6. Friston, K. (2010). The Free-Energy Principle: A Unified Brain Theory? *Nature Reviews Neuroscience*, 11(2), 127–138.
7. Clark, A. (2013). Whatever Next? Predictive Brains, Situated Agents, and the Future of Cognitive Science. *Behavioral and Brain Sciences*, 36(3), 181–204.
8. Hohwy, J. (2013). *The Predictive Mind*. Oxford University Press.
9. Liu, Y.-Y., Slotine, J.-J., & Barabási, A.-L. (2011). Controllability of Complex Networks. *Nature*, 473(7346), 167–173.
10. Gu, S., Pasqualetti, F., Cieslak, M., et al. (2015). Controllability of Structural Brain Networks. *Nature Communications*, 6, 8414.

11. Karrer, B., Newman, M. E. J., & Zdeborová, L. (2014). Percolation on Sparse Networks. *Physical Review Letters*, 113, 208702.
12. Kelso, J. A. S. (1995). *Dynamic Patterns: The Self-Organization of Brain and Behavior*. MIT Press.
13. Mitchell, M. (2009). *Complexity: A Guided Tour*. Oxford University Press.
14. Russell, S. (2019). *Human Compatible: Artificial Intelligence and the Problem of Control*. Viking Press.
15. Rahwan, I., et al. (2019). Machine Behaviour. *Nature*, 568(7753), 477–486.
16. Gabriel, I. (2020). Artificial Intelligence, Values, and Alignment. *Minds and Machines*, 30(3), 411–437.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.