

Article

Not peer-reviewed version

The Informational Coherence Index: A Metric for Evaluating Integration in Networks of Artificial Intelligence Models

[Henry Matuchaki](#) *

Posted Date: 12 February 2026

doi: 10.20944/preprints202502.2063.v2

Keywords: informational coherence index; AI model networks; network integration; entropy; Gaussian coupling; resonance; distributed AI architecture; gradient-based optimization



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.

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.

Article

The Informational Coherence Index: A Metric for Evaluating Integration in Networks of Artificial Intelligence Models

Henry Matuchaki

Independent Researcher, Brazil; henrymatuchaki@gmail.com

Abstract

We introduce the Informational Coherence Index (I_{coer}), a bounded metric for quantifying the degree of integration and alignment among interconnected artificial intelligence (AI) models. The index combines four interpretable components: normalized processing capacity, a Gaussian informational coupling function that decays with inter-model distance, an entropy-based weight reflecting output uncertainty, and a Lorentzian resonance factor capturing synchronicity. We prove that $I_{\text{coer}} \in [0, 1]$ and is monotonically decreasing in both informational distance and entropy. We derive closed-form gradients for optimizing network coherence via gradient ascent and demonstrate convergence on networks of up to 100 models. A pairwise extension of the index is also proposed for settings where inter-model distances are defined over embedding spaces. Simulation experiments confirm that the metric exhibits proportional sensitivity to parameter variations, correctly identifies informational bottlenecks, and scales to large networks. We discuss the relationship between I_{coer} and established measures such as agreement rate and mutual information, and outline directions for empirical validation on real AI systems.

Keywords: informational coherence index; AI model networks; network integration; entropy; Gaussian coupling; resonance; distributed AI architecture; gradient-based optimization

1. Introduction

The deployment of artificial intelligence (AI) increasingly relies on networks of heterogeneous models that collaborate to perform complex tasks. Model ensembles [7], multi-agent systems [16], and federated learning architectures [10] all involve collections of models whose outputs must be integrated, compared, or coordinated. As these networks grow in size and diversity, the need for principled metrics to evaluate the quality of their integration becomes critical.

Existing evaluation approaches tend to focus on individual aspects of model interaction. Agreement rate measures the fraction of concordant predictions between pairs of models. Cosine similarity quantifies the alignment of embedding vectors. Mutual information captures shared statistical structure between outputs. Cross-entropy loss evaluates predictive divergence. While each of these metrics is valuable in isolation, none provides a unified measure of *coherence*—the degree to which a network of models functions as an integrated, aligned system.

In this work, we propose the *Informational Coherence Index* (I_{coer}), a composite metric designed to capture multiple dimensions of network integration simultaneously. The index combines four interpretable factors: the processing capacity of each model, the informational distance between models, the entropy of model outputs, and the resonance (synchronicity) among models. By construction, I_{coer} is bounded in $[0, 1]$, monotonically decreasing in both distance and entropy, and amenable to gradient-based optimization.

The formulation draws on well-established mathematical tools. The coupling between models is modeled as a Gaussian function of informational distance, providing smooth decay without singularities. Output uncertainty is captured through Shannon entropy weighted by an exponential factor.

Synchronicity is measured via a Lorentzian resonance function centered on a reference frequency. The combination of these terms produces a metric with clear physical and information-theoretic interpretability.

We make the following contributions:

1. We formally define I_{coer} and prove its mathematical properties (boundedness, monotonicity, differentiability).
2. We derive closed-form gradients enabling gradient-based optimization of network coherence.
3. We propose a pairwise extension for settings where inter-model distances are defined in embedding spaces.
4. We validate the metric through simulation experiments on networks of 5 to 100 models, demonstrating proportional sensitivity, convergence of optimization, and scalability.
5. We discuss connections to existing metrics and outline directions for real-world validation.

The remainder of this paper is organized as follows. Section 2 reviews related work. Section 3 presents the formal definition of I_{coer} and its properties. Section 4 derives the optimization framework. Section 5 reports simulation results. Section 6 discusses implications and limitations. Section 7 concludes.

2. Related Work

Model Ensembles.

Ensemble methods combine the predictions of multiple models to improve accuracy and robustness. Lakshminarayanan et al. [7] introduced deep ensembles as a simple and effective approach to uncertainty estimation, while Fort et al. [3] analyzed the role of diversity in ensemble performance. BatchEnsemble [15] proposed parameter-efficient ensembling for large-scale networks. These works typically evaluate ensembles through predictive accuracy and calibration, but do not provide a unified coherence metric for the network as a whole.

Federated Learning.

Federated learning [10] trains models across distributed nodes while preserving data privacy. The FedAvg algorithm and its variants address the challenge of aggregating heterogeneous local models into a coherent global model. Karimireddy et al. [5] addressed client drift, while Li et al. [9] studied the convergence of federated optimization under heterogeneous data distributions. These works focus on convergence of the training process rather than on measuring the coherence of the resulting model network.

Information-Theoretic Measures.

Shannon entropy [11] provides the foundational measure of uncertainty in model outputs. Mutual information has been used to quantify shared information between model representations [13]. KL divergence and cross-entropy are standard tools for measuring distributional differences between models. Smith et al. [12] proposed methods for measuring the alignment of large language models. The Informational Coherence Index builds upon these foundations by integrating entropy into a multi-factor metric that also accounts for distance, capacity, and resonance.

Network Science and AI.

The analysis of complex networks [1] provides tools for studying the structure and dynamics of interconnected systems. Graph-based representations of model interactions have been applied to multi-agent systems and neural architecture search. The present work draws on the network perspective by treating AI models as nodes in a graph, where coherence is determined by the aggregate quality of pairwise interactions.

Physical Analogies in Machine Learning.

The use of physical analogies in machine learning has a rich history. Energy-based models [8], Boltzmann machines [4], and simulated annealing [6] all draw on thermodynamic concepts. The temperature parameter in softmax distributions and in sampling strategies (e.g., in language models) is directly inspired by statistical mechanics. Our use of an exponential entropy weight $e^{-\beta S_i}$ follows this tradition, with β serving as a tunable hyperparameter controlling the sensitivity to output uncertainty.

3. The Informational Coherence Index

We consider a network of N AI models, each characterized by its processing capacity, the informational distance to a reference point (or to other models), the entropy of its outputs, and its operating frequency. The Informational Coherence Index (I_{coer}) aggregates these factors into a single bounded metric.

3.1. Definition

Definition 1 (Informational Coherence Index). *Let $\{M_1, M_2, \dots, M_N\}$ be a network of N AI models. For each model M_i , let $C_i > 0$ denote its processing capacity, $r_i \geq 0$ its informational distance, $S_i \geq 0$ its output entropy, and $\Gamma_i \in (0, 1]$ its resonance factor. The Informational Coherence Index is defined as:*

$$I_{\text{coer}} = \frac{1}{N} \sum_{i=1}^N \hat{C}_i \cdot \varphi(r_i; \sigma) \cdot e^{-\beta S_i} \cdot \Gamma_i \quad (1)$$

where $\hat{C}_i = C_i / \max_j C_j$ is the normalized capacity, $\varphi(r; \sigma) = \exp(-r^2/2\sigma^2)$ is the Gaussian coupling function with scale parameter $\sigma > 0$, and $\beta > 0$ is a temperature hyperparameter.

Each factor captures a distinct dimension of network integration:

- $\hat{C}_i \in (0, 1]$: Models with higher processing capacity contribute more to coherence.
 - $\varphi(r_i; \sigma) \in (0, 1]$: Models closer in informational space contribute more; coupling decays smoothly with distance.
 - $e^{-\beta S_i} \in (0, 1]$: Models with lower output entropy (less uncertainty) contribute more.
 - $\Gamma_i \in (0, 1]$: Models better synchronized with the network's reference frequency contribute more.
- Since each factor lies in $(0, 1]$ and the sum is divided by N , we have $I_{\text{coer}} \in (0, 1]$.

3.2. Processing Capacity

The processing capacity C_i reflects the computational complexity and scale of model M_i . For models characterized by a parameter count N_{param} , we define:

$$C_i = \alpha \cdot \log(N_{\text{param}}^{(i)} + 1) \quad (2)$$

where $\alpha > 0$ is a scaling factor adjustable to the computational environment. The logarithmic transformation ensures that capacity grows sublinearly with model size, reflecting the diminishing returns of adding parameters. The normalization $\hat{C}_i = C_i / \max_j C_j$ ensures that the largest model in the network has $\hat{C}_i = 1$.

In practice, C_i can also incorporate other measures of computational power, such as FLOPs (floating-point operations per second) or inference throughput.

3.3. Gaussian Informational Coupling

The coupling function $\varphi(r; \sigma)$ measures the strength of informational interaction as a function of distance r :

$$\varphi(r; \sigma) = \exp\left(-\frac{r^2}{2\sigma^2}\right) \quad (3)$$

This Gaussian form is chosen for several reasons:

1. **Correct monotonicity:** φ is strictly decreasing for $r > 0$, ensuring that closer models contribute more to coherence.
2. **Smoothness:** φ is infinitely differentiable, enabling gradient-based optimization.
3. **No singularities:** Unlike power-law functions such as r^{-p} , the Gaussian is well-defined at $r = 0$.
4. **Interpretable scale:** The parameter σ has a clear interpretation as the characteristic distance at which coupling decays to $e^{-1/2} \approx 0.607$.

The informational distance r_i can be defined in multiple ways depending on the application:

- *Embedding distance:* $r_i = \|e_i - e_{\text{ref}}\|_2$, where e_i is the model's output embedding and e_{ref} is a reference embedding.
- *Data dissimilarity:* Measured by the divergence between the training data distributions.
- *Architectural distance:* Based on structural differences in model architectures.
- *Latency-based distance:* Reflecting communication delays in distributed systems.

Property 1 (Coupling Decay). For all $\sigma > 0$:

1. $\varphi(0; \sigma) = 1$ (maximum coupling at zero distance).
2. $\frac{\partial \varphi}{\partial r} = -\frac{r}{\sigma^2} \varphi(r; \sigma) < 0$ for $r > 0$ (strictly decreasing).
3. $\lim_{r \rightarrow \infty} \varphi(r; \sigma) = 0$ (vanishing coupling at infinite distance).

3.4. Entropy-Based Weight

The output entropy of model M_i is measured using Shannon entropy [11]:

$$S_i = - \sum_j p_{ij} \log p_{ij} \quad (4)$$

where p_{ij} is the probability assigned to the j -th output class (or token) by model M_i . High entropy indicates high uncertainty in the model's predictions; low entropy indicates confident, concentrated outputs.

The entropic weight is defined as:

$$w_i^{(S)} = e^{-\beta S_i} \quad (5)$$

where $\beta > 0$ is a temperature hyperparameter. This exponential weighting has the following properties:

- When β is small, all models contribute approximately equally regardless of entropy.
- When β is large, only models with very low entropy contribute significantly.
- The functional form $e^{-\beta S}$ is standard in statistical mechanics and maximum entropy methods, and provides a smooth, differentiable weighting.

We emphasize that β is treated as a tunable hyperparameter with no physical units. While the notation is borrowed from statistical mechanics, where $\beta = 1/(k_B T)$, no physical temperature or Boltzmann constant is implied. The parameter simply controls the trade-off between including all models (low β) and prioritizing low-entropy models (high β).

3.5. Harmonic Resonance Factor

The resonance factor Γ_i captures the degree of synchronicity between model M_i and the network's reference behavior. We model this using a Lorentzian (Cauchy) function:

$$\Gamma_i = \frac{1}{1 + (\omega_i - \omega_0)^2} \quad (6)$$

where ω_i is the operating frequency (or processing rate) of model M_i and ω_0 is the network's reference frequency.

This function achieves its maximum value of 1 when $\omega_i = \omega_0$ (perfect resonance) and decays symmetrically as the model's frequency deviates from the reference. The Lorentzian form is a natural

choice for resonance phenomena, appearing in the spectral line shapes of oscillating systems and in the frequency response of damped harmonic oscillators.

In practice, ω_i can represent the model's inference speed, update frequency in online learning, or response latency. The reference frequency ω_0 can be set to the network mean:

$$\omega_0 = \frac{1}{N} \sum_{i=1}^N \omega_i \quad (7)$$

or to a target frequency specified by the system designer.

3.6. Formal Properties

We establish the key mathematical properties of I_{coer} .

Proposition 1 (Boundedness). *For any network configuration, $0 < I_{\text{coer}} \leq 1$.*

Proof. Each factor satisfies: $\hat{C}_i \in (0, 1]$, $\varphi(r_i; \sigma) \in (0, 1]$, $e^{-\beta S_i} \in (0, 1]$ (for $S_i \geq 0$, $\beta > 0$), and $\Gamma_i \in (0, 1]$. Therefore, each summand in (1) lies in $(0, 1]$, and $I_{\text{coer}} = \frac{1}{N} \sum_{i=1}^N (\cdot) \in (0, 1]$. The upper bound is achieved when $\hat{C}_i = 1$, $r_i = 0$, $S_i = 0$, and $\Gamma_i = 1$ for all i . $\square$

Proposition 2 (Monotonicity in Distance). *For all i and for $r_i > 0$:*

$$\frac{\partial I_{\text{coer}}}{\partial r_i} = -\frac{r_i}{\sigma^2} \cdot \frac{\hat{C}_i \cdot \varphi(r_i; \sigma) \cdot e^{-\beta S_i} \cdot \Gamma_i}{N} < 0 \quad (8)$$

That is, I_{coer} is strictly decreasing in each r_i .

Proposition 3 (Monotonicity in Entropy). *For all i and $\beta > 0$:*

$$\frac{\partial I_{\text{coer}}}{\partial S_i} = -\beta \cdot \frac{\hat{C}_i \cdot \varphi(r_i; \sigma) \cdot e^{-\beta S_i} \cdot \Gamma_i}{N} < 0 \quad (9)$$

That is, I_{coer} is strictly decreasing in each S_i .

These properties confirm that I_{coer} behaves as expected for a coherence measure: the index increases when models are closer (lower r_i), more certain (lower S_i), and better synchronized (higher Γ_i).

3.7. Pairwise Extension

The formulation in (1) measures each model's coherence with respect to a reference point. In many settings, it is more natural to define coherence over all *pairs* of models. We propose the following pairwise extension:

Definition 2 (Pairwise Informational Coherence Index).

$$I_{\text{coer}}^{(pw)} = \frac{2}{N(N-1)} \sum_{i < j} C_{ij} \cdot \varphi(r_{ij}; \sigma) \cdot e^{-\beta S_{ij}} \cdot \Gamma_{ij} \quad (10)$$

where:

- $C_{ij} = \min(C_i, C_j) / \max(C_i, C_j) \in (0, 1]$ (capacity compatibility),
- $r_{ij} = \|e_i - e_j\|_2$ (Euclidean distance in embedding space),
- $S_{ij} = (S_i + S_j) / 2$ (mean pairwise entropy),
- $\Gamma_{ij} = 1 / (1 + (\omega_i - \omega_j)^2)$ (pairwise resonance).

The pairwise index inherits the boundedness and monotonicity properties of the pointwise version. Its computational cost is $O(N^2)$, which is acceptable for networks of moderate size (up to $\sim 10^3$ models). For larger networks, approximate methods based on random sampling of pairs or locality-sensitive hashing can be employed.

4. Optimization

A key advantage of I_{coer} is that its differentiability enables gradient-based optimization of network coherence. Given a network of models with fixed capacities C_i and resonance factors Γ_i , we seek to adjust the informational distances r_i and entropies S_i to maximize I_{coer} .

4.1. Analytical Gradients

From Propositions 2 and 3, the gradients of I_{coer} with respect to r_i and S_i are:

$$\frac{\partial I_{\text{coer}}}{\partial r_i} = -\frac{r_i}{\sigma^2 N} \hat{C}_i \varphi(r_i; \sigma) e^{-\beta S_i} \Gamma_i \quad (11)$$

$$\frac{\partial I_{\text{coer}}}{\partial S_i} = -\frac{\beta}{N} \hat{C}_i \varphi(r_i; \sigma) e^{-\beta S_i} \Gamma_i \quad (12)$$

Both gradients are strictly negative (for $r_i > 0$, $\beta > 0$), confirming that the objective is maximized by decreasing distances and entropies. The gradient ascent updates are:

$$r_i^{(t+1)} = r_i^{(t)} + \eta_r \cdot \frac{\partial I_{\text{coer}}}{\partial r_i} \quad (13)$$

$$S_i^{(t+1)} = S_i^{(t)} + \eta_S \cdot \frac{\partial I_{\text{coer}}}{\partial S_i} \quad (14)$$

where $\eta_r > 0$ and $\eta_S > 0$ are learning rates. Since the gradients are negative, the updates decrease r_i and S_i , driving the network toward higher coherence.

4.2. Diversity Regularization

Without regularization, the optimization drives all $r_i \rightarrow 0$ and $S_i \rightarrow 0$, collapsing the network to a degenerate state where all models are identical. To maintain structural diversity while maximizing coherence, we introduce a regularization term:

$$I_{\text{coer}}^{(\text{reg})} = I_{\text{coer}} + \lambda \cdot \text{std}(\mathbf{r}) \quad (15)$$

where $\text{std}(\mathbf{r})$ is the standard deviation of the distance vector and $\lambda \geq 0$ controls the strength of the diversity penalty. This encourages the network to maintain heterogeneity in informational distances while still optimizing for coherence.

4.3. Algorithm

The complete optimization procedure is presented in Algorithm 1.

Algorithm 1 Gradient Ascent Optimization of I_{coer}

```
Require:  $C_i, r_i^{(0)}, S_i^{(0)}, \Gamma_i, \beta, \sigma, \eta_r, \eta_S, \lambda, T_{\text{max}}, \epsilon$ 
1: Compute  $\hat{C}_i = C_i / \max_j C_j$ 
2: for  $t = 0, 1, \dots, T_{\text{max}}$  do
3:   Compute  $I_{\text{coer}}^{(t)}$  via (1)
4:   Compute  $\nabla_{r_i} I_{\text{coer}}^{(t)}$  via (11)
5:   Compute  $\nabla_{S_i} I_{\text{coer}}^{(t)}$  via (12)
6:   if  $\lambda > 0$  then
7:     Add diversity gradient:  $\nabla_{r_i} \leftarrow \nabla_{r_i} + \lambda \cdot \frac{r_i - \bar{r}}{N \cdot \text{std}(\mathbf{r})}$ 
8:   end if
9:   Update:  $r_i^{(t+1)} = \text{clip}(r_i^{(t)} + \eta_r \nabla_{r_i}, 0.01, r_{\text{max}})$ 
10:  Update:  $S_i^{(t+1)} = \text{clip}(S_i^{(t)} + \eta_S \nabla_{S_i}, 0.01, S_{\text{max}})$ 
11:  if  $|I_{\text{coer}}^{(t)} - I_{\text{coer}}^{(t-1)}| < \epsilon \cdot I_{\text{coer}}^{(t-1)}$  and  $t > 10$  then
12:    break {Convergence}
13:  end if
14: end for
15: return  $r_i^{(t)}, S_i^{(t)}, I_{\text{coer}}^{(t)}$ 
```

The clip operation enforces physical constraints: distances and entropies must remain positive and bounded. The convergence criterion halts optimization when the relative change in I_{coer} falls below a threshold ϵ .

5. Experiments

We validate I_{coer} through a series of simulation experiments designed to test its mathematical properties, sensitivity, optimization dynamics, and scalability.

5.1. Experimental Setup

All experiments use the following parameter ranges unless otherwise noted:

- Processing capacities $C_i \sim \mathcal{U}(80, 120)$.
- Informational distances $r_i \sim \mathcal{U}(0.5, 5.0)$.
- Output entropies $S_i \sim \mathcal{U}(0.3, 0.9)$.
- Operating frequencies $\omega_i \sim \mathcal{U}(0.5, 1.5)$, with $\omega_0 = 1.0$.
- Temperature $\beta = 1.0$, coupling scale $\sigma = 1.5$.

Resonance factors are computed analytically as $\Gamma_i = 1 / (1 + (\omega_i - \omega_0)^2)$. All results are averaged over 10 random seeds where applicable.

5.2. Five-Model Network

We first examine a small network of five models with fixed parameters to verify the contribution distribution. Table 1 shows the per-model breakdown.

Table 1. Per-model contributions to I_{coer} in a five-model network ($\sigma = 1.5, \beta = 1.0$).

Model	C_i	r_i	S_i	ω_i	Γ_i	Contribution (%)
M_0	100	1.0	0.5	1.00	1.000	24.7
M_1	80	1.2	0.7	1.10	0.990	14.5
M_2	120	1.4	0.3	0.90	0.990	28.9
M_3	90	1.6	0.6	1.20	0.962	13.7
M_4	110	1.8	0.4	0.95	0.998	18.2
I_{coer}						0.328

The contributions are distributed across all models, with model M_2 contributing the most (28.9%) due to its combination of high capacity ($C_2 = 120$) and low entropy ($S_2 = 0.3$). Importantly, no single model dominates the index, confirming that the metric captures collective behavior rather than being determined by a single outlier.

5.3. Sensitivity Analysis

We evaluate the sensitivity of I_{coer} to $\pm 10\%$ perturbations of each parameter. Table 2 reports the results.

Table 2. Sensitivity of I_{coer} to $\pm 10\%$ parameter perturbations.

Parameter	+10% change in I_{coer}	−10% change in I_{coer}
r_i (distance)	−8.7%	+8.7%
S_i (entropy)	−4.5%	+4.8%
Γ_i (resonance)	+10.0%	−10.0%

The sensitivity is proportional and well-behaved: a 10% change in any parameter produces a change in I_{coer} of at most 10%. The index is most sensitive to informational distance, followed by resonance and entropy. This proportional sensitivity stands in contrast to formulations based on high-order power laws, which can produce sensitivity of millions of percent for small parameter changes.

5.4. Optimization Convergence

We apply Algorithm 1 to a network of 20 models with random initial parameters. The optimization uses learning rates $\eta_r = 0.1$ and $\eta_S = 0.05$, with coupling scale $\sigma = 2.0$. Table 3 shows the trajectory.

Table 3. Optimization trajectory for a 20-model network.

Iteration	I_{coer}	$\bar{r}$	$\bar{S}$
0	0.2165	2.547	0.594
10	0.2188	2.542	0.589
25	0.2222	2.535	0.580
50	0.2282	2.523	0.566
100	0.2413	2.498	0.537
200	0.2725	2.446	0.474
500	0.3652	2.346	0.339

The optimization produces a monotonic increase in I_{coer} (+68.7% improvement), driven by decreasing mean distances and entropies. The convergence is smooth and consistent with the analytical gradient structure.

5.5. Scalability

To evaluate scalability, we apply the optimization procedure to networks of varying sizes. Table 4 reports the results.

Table 4. Scalability: optimization results across network sizes ($\sigma = 2.0$, 300 iterations).

N (models)	$I_{\text{coer}}^{(0)}$	$I_{\text{coer}}^{(\text{final})}$	Improvement (%)
5	0.3279	0.5413	65.1
20	0.2165	0.2969	37.2
50	0.1984	0.2237	12.8
100	0.1925	0.2049	6.4

As the network size increases, the initial I_{coer} decreases (more heterogeneous models) and the relative improvement from optimization decreases. This is expected: larger networks are harder to optimize due to the greater diversity of model configurations. Nevertheless, the optimization consistently improves coherence across all network sizes.

5.6. Correlation with Agreement Rate

To validate that I_{coer} captures a meaningful notion of network integration, we compute its correlation with agreement rate in a simulated setting. We generate 50 random networks with varying degrees of inter-model similarity (controlled by the spread of embedding vectors) and compute both $I_{\text{coer}}^{(\text{pw})}$ and a simulated agreement rate $a = e^{-d/3}$, where d is the embedding spread.

The Pearson correlation between $I_{\text{coer}}^{(\text{pw})}$ and agreement rate is $r = 0.785$ ($p < 0.001$), indicating a strong positive relationship. This confirms that I_{coer} captures a meaningful dimension of network coherence that aligns with the intuitive notion of model agreement.

6. Discussion

6.1. Relationship to Existing Metrics

I_{coer} is not a replacement for existing metrics but rather a complementary measure that captures multiple dimensions of network integration simultaneously. Table 5 summarizes the relationship.

Table 5. Relationship between I_{coer} and existing metrics.

Metric	Measures	Relationship to I_{coer}
Agreement Rate	Output concordance	Positive correlation ($r = 0.785$)
Cosine Similarity	Embedding alignment	Related to coupling term φ
Mutual Information	Shared structure	Related to entropy term $e^{-\beta S}$
Cross-Entropy Loss	Predictive divergence	Inversely related to coherence

Unlike these individual metrics, I_{coer} provides a single bounded score that jointly accounts for capacity, distance, entropy, and synchronicity.

6.2. Hyperparameter Selection

The index has three hyperparameters: σ (coupling scale), β (temperature), and λ (diversity regularization). We offer the following guidelines:

- σ should be set to the typical scale of informational distances in the network. A reasonable default is the median pairwise distance.
- β controls the trade-off between inclusivity and selectivity. Values in $[0.5, 2.0]$ provide a good balance for most applications.
- λ should be set to a small value (e.g., 0.01–0.1) when diversity preservation is desired, or to zero when maximum coherence is the sole objective.

6.3. Limitations

- Several limitations should be acknowledged:
1. **Simulated validation:** All experiments use synthetic data. Validation on real AI networks (e.g., federated learning systems, production ensembles) is necessary to confirm the practical utility of I_{coer} .
 2. **Distance definition:** The choice of informational distance r_i significantly affects the metric. Different definitions (embedding, architectural, latency-based) may yield different coherence assessments for the same network.
 3. **Pairwise scalability:** The pairwise extension has $O(N^2)$ complexity, which may be prohibitive for very large networks ($N > 10^4$).

4. **Static analysis:** The current formulation evaluates coherence at a single time point. Dynamic extensions for time-varying networks remain for future work.
5. **Fixed capacity and resonance:** The optimization treats C_i and Γ_i as fixed. In practice, both may evolve during training or deployment.

6.4. Future Directions

Several directions for future research are suggested:

- **Empirical validation:** Apply I_{coer} to real ensemble and federated learning systems, evaluating correlation with downstream task performance.
- **Dynamic extension:** Develop a time-dependent version $I_{\text{coer}}(t)$ for monitoring coherence during training.
- **Scalable pairwise computation:** Use sampling or locality-sensitive hashing to approximate $I_{\text{coer}}^{(\text{pw})}$ for large networks.
- **Integration with ML pipelines:** Incorporate I_{coer} as a regularization term in ensemble training objectives.
- **Cross-domain application:** Extend to scientific simulation networks, autonomous systems, and multi-modal AI.

7. Conclusion

We have introduced the Informational Coherence Index (I_{coer}), a bounded, differentiable metric for quantifying the integration quality of AI model networks. The index combines four interpretable components—capacity, coupling, entropy, and resonance—into a single score with provable mathematical properties: boundedness in $[0, 1]$, monotonicity in distance and entropy, and smooth differentiability enabling gradient-based optimization.

Through simulation experiments, we demonstrated that I_{coer} exhibits proportional sensitivity to parameter variations, converges under gradient ascent, scales to networks of 100 models, and correlates strongly with agreement rate. The pairwise extension provides a rigorous formulation for settings where inter-model distances are naturally defined.

While the current validation is simulation-based, the mathematical foundations and computational framework presented here provide a solid basis for future empirical applications. By offering a principled, unified measure of network coherence, I_{coer} contributes a new tool for the design, evaluation, and optimization of collaborative AI systems.

Appendix A Reference Implementation

The following Python code provides a reference implementation of I_{coer} and the gradient-based optimization procedure.

Listing 1: Core Icoer computation.

```

1 import numpy as np
2
3 def coupling_gaussian(r, sigma=1.5):
4     """Gaussian coupling:  $\phi(r) = \exp(-r^2 / 2\sigma^2)$ ."""
5     return np.exp(-r**2 / (2 * sigma**2))
6
7 def resonance(omega_i, omega_0):
8     """Lorentzian resonance:  $\Gamma = 1/(1+(\omega - \omega_0)^2)$ ."""
9     return 1.0 / (1.0 + (omega_i - omega_0)**2)
10
11 def icoer(C_i, r_i, S_i, Gamma_i, beta=1.0, sigma=1.5):
12     """Informational Coherence Index (Eq. 1)."""
13     N = len(C_i)
14     C_norm = C_i / np.max(C_i)

```

```

15     phi = coupling_gaussian(r_i, sigma)
16     w_entropy = np.exp(-beta * S_i)
17     return np.sum(C_norm * phi * w_entropy * Gamma_i) / N

```

Listing 2: Gradient-based optimization.

```

1 def grad_r(C_i, r_i, S_i, Gamma_i, beta, sigma):
2     """Gradient of Icoer w.r.t. r_i (Eq. 7)."""
3     N = len(C_i)
4     C_norm = C_i / np.max(C_i)
5     phi = coupling_gaussian(r_i, sigma)
6     w = np.exp(-beta * S_i)
7     return -r_i / (sigma**2) * C_norm * phi * w * Gamma_i / N
8
9 def grad_S(C_i, r_i, S_i, Gamma_i, beta, sigma):
10    """Gradient of Icoer w.r.t. S_i (Eq. 8)."""
11    N = len(C_i)
12    C_norm = C_i / np.max(C_i)
13    phi = coupling_gaussian(r_i, sigma)
14    w = np.exp(-beta * S_i)
15    return -beta * C_norm * phi * w * Gamma_i / N
16
17 def optimize(C_i, r_i, S_i, Gamma_i, beta=1.0,
18             sigma=1.5, lr_r=0.1, lr_S=0.05, iters=500):
19     """Gradient ascent to maximize Icoer."""
20     r, S = r_i.copy(), S_i.copy()
21     for t in range(iters):
22         r = np.clip(r + lr_r * grad_r(C_i, r, S, Gamma_i, beta, sigma),
23                     0.01, 10.0)
24         S = np.clip(S + lr_S * grad_S(C_i, r, S, Gamma_i, beta, sigma),
25                     0.01, 5.0)
26     return r, S, icoer(C_i, r, S, Gamma_i, beta, sigma)

```

Listing 3: Pairwise extension (Eq. 5).

```

1 def icoer_pairwise(C_i, embeddings, S_i, omega_i,
2                   beta=1.0, sigma=2.0):
3     """Pairwise Icoer over all model pairs."""
4     N = len(C_i)
5     total, n_pairs = 0.0, 0
6     for i in range(N):
7         for j in range(i+1, N):
8             C_ij = min(C_i[i], C_i[j]) / max(C_i[i], C_i[j])
9             r_ij = np.linalg.norm(embeddings[i] - embeddings[j])
10            phi = coupling_gaussian(r_ij, sigma)
11            S_ij = (S_i[i] + S_i[j]) / 2
12            G_ij = resonance(omega_i[i], omega_i[j])
13            total += C_ij * phi * np.exp(-beta * S_ij) * G_ij
14            n_pairs += 1
15    return total / n_pairs

```

References

1. A.-L. Barabási. *Linked: The New Science of Networks*. Perseus Publishing, 2002.
2. T. Brown, B. Mann, N. Ryder, et al. Language models are few-shot learners. In *NeurIPS*, 2020.

3. S. Fort, H. Hu, and B. Lakshminarayanan. Deep ensembles: A loss landscape perspective. *arXiv preprint arXiv:1912.02757*, 2019.
4. G. E. Hinton. Training products of experts by minimizing contrastive divergence. *Neural Computation*, 14(8):1771–1800, 2002.
5. S. P. Karimireddy, S. Kale, M. Mohri, S. J. Reddi, S. Stich, and A. T. Suresh. SCAFFOLD: Stochastic controlled averaging for federated learning. In *ICML*, 2020.
6. S. Kirkpatrick, C. D. Gelatt, and M. P. Vecchi. Optimization by simulated annealing. *Science*, 220(4598):671–680, 1983.
7. B. Lakshminarayanan, A. Pritzel, and C. Blundell. Simple and scalable predictive uncertainty estimation using deep ensembles. In *NeurIPS*, 2017.
8. Y. LeCun, S. Chopra, R. Hadsell, M. Ranzato, and F. Huang. A tutorial on energy-based learning. In *Predicting Structured Data*. MIT Press, 2006.
9. T. Li, A. K. Sahu, M. Zaheer, M. Sanjabi, A. Talwalkar, and V. Smith. Federated optimization in heterogeneous networks. In *MLSys*, 2020.
10. B. McMahan, E. Moore, D. Ramage, S. Hampson, and B. A. y Arcas. Communication-efficient learning of deep networks from decentralized data. In *AISTATS*, 2017.
11. C. E. Shannon. A mathematical theory of communication. *Bell System Technical Journal*, 27(3):379–423, 1948.
12. S. Smith, et al. Measuring the alignment of large language models. *arXiv preprint*, 2022.
13. N. Tishby and N. Zaslavsky. Deep learning and the information bottleneck principle. In *IEEE Information Theory Workshop*, 2015.
14. A. Vaswani, N. Shazeer, N. Parmar, et al. Attention is all you need. In *NeurIPS*, 2017.
15. Y. Wen, D. Tran, and J. Ba. BatchEnsemble: An alternative approach to efficient ensemble and lifelong learning. In *ICLR*, 2020.
16. M. Wooldridge. *An Introduction to MultiAgent Systems*. John Wiley & Sons, 2nd edition, 2009.

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.