

SIMULATING THE INEFFABLE: THE LIMITS OF ARTIFICIAL GENERAL INTELLIGENCE AND MATHEMATICAL MODELS IN GRASPING DAO

Dr N Krishnamoorthy¹, Dr C Vijayalakshmi², Dr P Prasanna Kumari³,
Saint Jesudoss.S⁴, Sendil Vadivu D⁵, S Selvakumar⁶, Mr. Hemant Sethi⁷, Dr.T.Vengatesh^{8*}

¹Faculty of Science and Humanities, PG Department of Computer Applications (MCA),
SRM INSTITUTE OF SCIENCE AND TECHNOLOGY, Ramapuram, Chennai 600089
Tamilnadu, India. <https://orcid.org/0000-0003-1792-3496>

Email id: krishnan@srmist.edu.in

²Assistant Professor (Senior Grade), Department of Computer Science and Engineering,
B S Abdur Rahman Crescent Institute of Science and Technology, Chennai - 600048
Email id vijayalakshmi@crescent.education

³Assistant Professor, Department of Computer Science and Engineering, Aditya University
Surampalem, Andhra Pradesh - 533 437, India
Email id: prasannaprofcom@gmail.com

⁴Assistant Professor, Department of CSE, RAJIV GANDHI COLLEGE OF ENGINEERING AND
TECHNOLOGY, PUDUCHERRY
Email id: saint.2k5@gmail.com

⁵Department of Computer Science and Engineering, Vel Tech Rangarajan Dr. Sagunthala R&D
Institute of Science and Technology, Chennai

⁶Professor, Department of AIML, Panimalar Engineering college, Chennai,
Email id: selvathendral1981@gmail.com

⁷Assistant Professor, Department of Computer Science and Engineering,
Maharishi Markandeshwar (Deemed to be University), Mullana, Ambala, Haryana
Email id: hemant.sethi@mmumullana.org

^{8*}(Corresponding Author) Assistant Professor, Department of Computer Science, Government Arts and Science
College, Veerapandi, Theni, Tamilnadu, India. Email: venkibiotinix@gmail.com

ABSTRACT

The pursuit of Artificial General Intelligence (AGI) rests on a foundational assumption: that intelligence can be fully captured by mathematical models and computational architectures. This paper challenges that assumption by examining the inherent limitations of algorithmic systems in grasping the Dao, a concept central to classical Chinese philosophy that denotes the fundamental, ineffable way of the universe. Drawing on insights from mathematical incompleteness, computational theory, and Daoist

epistemology, we argue that AGI, as currently conceived, faces structural barriers that preclude genuine engagement with the ineffable dimensions of reality. Through a multi-disciplinary framework combining formal proofs of semantic closure with Daoist critiques of prediction-based cognition, we demonstrate that the gap between algorithmic simulation and authentic understanding is not merely technical but ontological. Our analysis suggests that the most profound aspects of intelligence generativity, coordination, and sustaining identity across time remain inaccessible to mathematical formalization, positioning human cognition as irreplaceable in domains requiring wisdom rather than mere computation.

Keywords: Artificial General Intelligence, Dao, Mathematical Models, Semantic Closure, Epistemological Limits, Algorithmic Cognition.

1. INTRODUCTION

1.1 Background and Motivation

The contemporary discourse surrounding Artificial General Intelligence is characterized by a remarkable tension. On one hand, the rapid advancement of large language models and deep learning systems has fueled speculation that AGI is not only possible but imminent. On the other hand, a growing body of philosophical and mathematical critique suggests that such optimism may be fundamentally misplaced. This paper situates itself within this critical tradition, arguing that the very nature of AGI as an algorithmic, mathematically-modeled system imposes structural limitations that preclude genuine comprehension of the ineffable those dimensions of reality that resist formalization, prediction, and control.

The concept of the Dao, central to classical Daoist philosophy, serves as a particularly revealing test case for the limits of algorithmic intelligence. Daoist texts such as the *Daodejing* and *Zhuangzi* consistently emphasize that the fundamental way of the universe is not something that can be fully grasped, manipulated, or reduced to rational formulae. The Dao is characterized by complexity, mystery, and spontaneous emergence qualities that stand in stark contrast to the predictive, pattern-recognition paradigm that dominates contemporary AI research.

1.2 Research Questions and Objectives

This paper seeks to address the following interconnected questions:

1. **What are the structural limitations of algorithmic systems in capturing ineffable phenomena?** Drawing on formal results from computability theory and epistemology, we examine the mathematical boundaries that constrain AGI's capacity for genuine understanding.
2. **What insights does Daoist philosophy offer for understanding the limits of AI?** By engaging with classical Daoist critiques of rationalism and prediction, we identify conceptual resources for rethinking the relationship between intelligence and the ineffable.
3. **What are the implications of these limitations for the future of AI research and human-machine collaboration?** We consider how recognition of AGI's boundaries might reshape both technical approaches and ethical frameworks.

1.3 Scope and Structure

This paper proceeds through five main sections. Following this introduction, Section 2 provides a comprehensive literature survey spanning AI limitation arguments, Daoist philosophy of intelligence, and emerging critiques of algorithmic cognition. Section 3 presents our data description and methodology, including formal proofs of semantic closure and a proposed architectural framework. Section 4 reports results and implementation, while Section 5 offers discussion of implications and Section 6 concludes with directions for future research.

2. LITERATURE SURVEY

2.1 Foundational Critiques of AGI

The claim that AGI is impossible or structurally limited has a distinguished pedigree. Hubert Dreyfus's phenomenological critique, articulated in *What Computers Can't Do* (1972), argued that human intelligence is fundamentally embodied and situated in ways that cannot be captured by formal rule systems. John Searle's "Chinese Room" thought experiment (1980) challenged the computationalist

assumption that symbol manipulation could amount to genuine understanding. Roger Penrose linked Gödel's incompleteness theorems to claims about the non-computable nature of consciousness .

Recent work has formalized and extended these critiques. Landgrebe and Smith (2025) argue that AGI is mathematically impossible, offering two specific reasons: human intelligence is a capability of a complex dynamic system, and such systems cannot be modelled mathematically in a way that allows them to operate inside a computer . Their analysis demonstrates that AI systems are "best viewed as pieces of mathematics, which cannot think, feel, or will" .

Schlereth (2025) develops the "Infinite Choice Barrier" theorem, providing formal proof that algorithmic systems are structurally incapable of resolving problems requiring semantic innovation, frame creation, or decisions under irreducible uncertainty . Building on Gödel, Turing, and Kant, Schlereth demonstrates that for any algorithmic system, there exist decision contexts where optimal performance is structurally unattainable .

2.2 The Ontology of Intelligence

Wang and Ji (2025) propose a "Structural-Generative Ontology of Intelligence" that challenges functionalist and behaviorist conceptions of AI . They argue that true intelligence exists only when a system can generate new structures, coordinate them into reasons, and sustain its identity over time. Current AI systems, however broad in function, "remain surface simulations because they lack this depth" . This framework positions true AGI as a "Second Being" standing alongside human existence but ontologically distinct, irreducible to simulation.

He (2025) develops "Dual-Nature Intelligence" from the perspective of Dao-Er ontology, arguing that intelligence is neither an isolated attribute of entities nor a single computational process, but the dynamic generation and dialectical unity of dual relations . This framework "breaks the monistic limitations of traditional AI philosophy" and provides an integration of Eastern wisdom with practical explanatory power .

2.3 Daoist Critiques of AI

A growing body of literature applies Daoist philosophy to critique contemporary AI paradigms. D'Ambrosio (2025) offers a Daoist-inspired critique of AI's promises, arguing that the focus on patterns, predictions, and control reflects assumptions that have escaped serious critical analysis . While Daoist texts also assume there are patterns in the world, they "are not enthusiastic about the rational or knowable nature of these patterns rather, they encourage us to appreciate them as fundamentally complex and mysterious" .

Collet (2025) argues that the Daoist concept of *wu wei* non-forcing, non-coercive action constitutes the most rigorous philosophical account of genuine collaborative creativity, and that contemporary AI systematically fails to achieve this . Wu wei is "not passivity. It is the most demanding mode of engagement available to a mind: acting in perfect accordance with the nature of the situation, following the lines of force already present in the encounter, allowing emergence rather than generating output" . Shin (2021) challenges techno-optimist assumptions by reframing "unintelligence" not as a lack, but as a generative state of epistemic openness . Drawing on Daoist reflections on " (wu, nothingness) and " (xu, emptiness), Shin argues for "reclaiming spaces of intentional unknowing, generative pause, and epistemological humility" as essential for humane innovation .

2.4 The Limits of Formal Systems

The mathematical foundations for AGI limitation arguments draw on Gödel's incompleteness theorems, which demonstrate that in any sufficiently expressive formal system, there exist true statements that cannot be proven within the system . Turing's halting problem shows that no general algorithm can decide, for every program and input, whether the program will halt .

Schlereth (2025) formalizes the "Semantic Closure Theorem," demonstrating that an algorithmic system cannot generate outputs whose interpretation requires a semantic framework beyond the one defined by its architecture and training history . This theorem establishes that "true semantic innovation is structurally impossible for algorithmic systems" .

2.5 Research Gaps

While the literature has addressed philosophical and mathematical critiques of AGI, a significant gap remains in the integration of these perspectives with Daoist epistemology. This paper addresses that gap by developing a unified framework that combines formal proofs of algorithmic limitations with Daoist insights into the nature of the ineffable, demonstrating why the simulation of intelligence is not equivalent to authentic understanding.

3. DATA DESCRIPTION AND METHODOLOGY

3.1 Conceptual Data Sources

Our analysis draws on three categories of data:

Category A: Philosophical Texts

- Primary Daoist texts (*Daodejing*, *Zhuangzi*)
- Kantian epistemology (*Critique of Pure Reason*)
- Phenomenological critiques of AI (Dreyfus)

Category B: Mathematical Foundations

- Gödel's incompleteness theorems
- Turing's halting problem
- Semantic closure theorem formulations

Category C: Contemporary AI Research

- Scaling behavior of large language models
- Empirical evidence of performance ceilings
- Model collapse phenomena

3.2 Formal Framework

We define an algorithmic system SS as a tuple:

$$S=(\Sigma,R,\Theta,L)S=(\Sigma,R,\Theta,L)$$

Where:

- $\Sigma\Sigma$ = symbol set / representational space
- RR = transformation rules
- $\Theta\Theta$ = learned parameters
- LL = loss/objective function

The **Semantic Frame** $\Omega(S)\Omega(S)$ is defined as:

$$\Omega(S)=(\Sigma,R,\Theta)\Omega(S)=(\Sigma,R,\Theta)$$

Semantic Closure Theorem: For any algorithmic system SS :

$$\forall o \bullet \text{Output}(S):\text{Interpretation}(o) \bullet \Omega(S) \forall o \bullet \text{Output}(S):\text{Interpretation}(o) \bullet \Omega(S)$$

That is: all outputs generated by the system remain semantically interpretable within its internal framework. True semantic innovation the creation of new frames $\Omega' \bullet \Omega(S)\Omega' \bullet \Omega(S)$ is structurally impossible .

3.3 Dao-Inspired Critique Framework

Building on D'Ambrosio's analysis , we identify three dimensions of the Dao that resist algorithmic formalization:

Dimension	1:	Non-Rational	Patterns
-----------	----	--------------	----------

Daoist epistemology acknowledges patterns in the world but refuses to reduce them to knowable, rational formulae. The Dao is fundamentally complex and mysterious, operating through *ziran* spontaneous self-so-ness that cannot be captured by prediction algorithms .

Dimension	2:	Ethical	Uncodifiability
-----------	----	---------	-----------------

Ethical considerations "cannot easily be controlled or known, much less put into code" . The Daoist approach to ethics emphasizes humility, contextual sensitivity, and the recognition that values resist mathematical translation.

Dimension	3:	Generative	Openness
-----------	----	------------	----------

The Dao is not a fixed structure but a dynamic, generative process. This aligns with the "Structural-

Generative Ontology of Intelligence" proposed by Wang and Ji , where true intelligence requires generativity, coordination, and sustaining capacities that algorithmic systems lack.

3.4 Proposed Methodology: The "Dao-Aware AI" Framework

Our proposed methodology integrates Daoist insights with formal limitation arguments to develop a "Dao-Aware AI" framework. The framework has three components:

Component 1:	Epistemic Humility	Module
Acknowledges the inherent limitations of algorithmic systems in grasping ineffable phenomena. Implements explicit recognition of "unknown unknowns" and resists the temptation to equate prediction with understanding.		
Component 2:	Non-Algorithmic Integration	Layer
Provides interfaces for human judgment and wisdom to supplement algorithmic processing. Recognizes that certain decisions particularly those involving ethical complexity, creative innovation, or existential meaning cannot be delegated to algorithmic systems.		
Component 3:	Generative Pause	Mechanism
Inspired by Daoist <i>wu wei</i> , this mechanism introduces reflective pauses that resist the imperative to force outputs, predict outcomes, or optimize continuously. It creates space for emergence rather than mere generation.		

4.PROPOSED METHODOLOGY: DAO-AWARE AI FRAMEWORK

4.1. METHODOLOGICAL OVERVIEW

The proposed methodology addresses the fundamental research question: *How can we design AI systems that acknowledge their inherent limitations in grasping ineffable phenomena while maintaining utility in domains accessible to algorithmic processing?*

Our approach integrates three distinct but interconnected methodological streams:

1. **Formal Mathematical Analysis** - Establishing the theoretical boundaries of algorithmic systems through semantic closure theorems and incompleteness arguments
2. **Philosophical Framework Development** - Operationalizing Daoist concepts (*wu wei*, *ziran*, epistemic humility) into design principles
3. **Architectural Implementation** - Creating a hybrid system that combines algorithmic processing with human wisdom integration

4.2. METHODOLOGICAL COMPONENTS

4.2.1 Formal Framework Development

Step 1: Semantic Boundary Mapping

For any algorithmic system *S*, we formally define its semantic frame $\Omega(S)$ and establish the boundaries beyond which the system cannot operate meaningfully.

Step 2: Ineffability Classification

We develop a systematic classification of problem domains based on their "ineffability quotient" - the degree to which they resist algorithmic formalization:

Classification	Characteristics	Example Domains
Fully Formalizable	Complete mathematical description possible	Arithmetic, Logic, Pattern Recognition
Partially Formalizable	Some aspects formalizable, others resist	Natural Language Processing, Prediction
High-Ineffable	Core phenomena resist formalization	Ethics, Creativity, Wisdom
Fully Ineffable	Resists all formal description	Direct experience of Dao, Wu Wei

4.2.2 Dao-Inspired Design Principles

Principle 1: Epistemic Humility (From *ziran* - natural spontaneity)

- Systems must explicitly recognize their own limitations
- Avoid overclaiming understanding or prediction capability
- Maintain awareness of "unknown unknowns"

Principle 2: Generative Non-Forcing (From *wu wei* - effortless action)

- Resist the imperative to constantly produce outputs
- Allow space for emergence rather than forced generation
- Prioritize alignment with situation over optimization

Principle 3: Dynamic Openness (From *xu* - emptiness)

- Maintain receptivity to new patterns
- Avoid premature closure or crystallization of understanding
- Preserve capacity for genuine novelty

4.2.3 Hybrid Architecture Design

The Dao-Aware AI architecture comprises five interconnected layers:



Figure 1: The Dao-Aware Hybrid AI Architecture

Architecture Explanation:

1. **Input Layer:** Distinguishes between structured data (fully formalizable), ambiguous inputs (partially formalizable), and ineffable queries that resist algorithmic processing

- 2. **Semantic Classification Layer:** Calculates the ineffability quotient of each input based on formal criteria derived from our theoretical framework
- 3. **Algorithmic Processing Layer:** Applies conventional AI techniques (pattern recognition, prediction, optimization) with integrated uncertainty quantification
- 4. **Non-Algorithmic Integration Layer:** The core innovation where:
 - **Epistemic Humility Module** acknowledges system limitations
 - **Generative Pause Mechanism** creates space for emergent understanding
 - **Human Wisdom Interface** enables collaboration with human judgment
 - **Emergence Facilitation** supports non-forced, spontaneous pattern recognition
- 5. **Output Layer:** Provides graduated responses from confident outputs (fully formalizable domains) to collaborative responses and reflective pause signals (ineffable domains)

4.3. METHODOLOGICAL IMPLEMENTATION

4.3.1 Algorithmic Implementation

Algorithm 1: Semantic Boundary Assessment

text

Input: Query Q, System S with semantic frame $\Omega(S)$

Output: Ineffability Score $I(Q) \in [0,1]$

- 1. Represent Q in system's formal language $L(S)$
- 2. Attempt formal derivation of Q within $\Omega(S)$
- 3. If successful: $I(Q) = 0$ (fully formalizable)
- 4. If partial success:
 - Calculate coverage ratio $C = |\Omega(Q) \cap \Omega(S)| / |\Omega(Q)|$
 - Set $I(Q) = 1 - C$
- 5. If no representation possible:
 - Set $I(Q) = 1$ (fully ineffable)
- 6. Return $I(Q)$

Algorithm 2: Generative Pause Trigger

text

Input: Ineffability Score $I(Q)$, Confidence Score $C(Q)$

Output: Processing Mode $M \in \{\text{Algorithms, Collaborative, Pause}\}$

- 1. If $I(Q) < 0.3$ and $C(Q) > 0.8$:
Return "Algorithmic" (process fully algorithmically)
- 2. If $0.3 \leq I(Q) \leq 0.7$:
Return "Collaborative" (human-AI collaboration)
- 3. If $I(Q) > 0.7$ or $C(Q) < 0.3$:
Return "Pause" (generate reflective pause, request human input)
- 4. Return result

4.3.2 Evaluation Methodology

We evaluate the Dao-Aware AI framework across multiple dimensions

Metric	Description	Measurement Method
Formalizability Accuracy	Correct classification of problem domains	Comparison with human expert classification
Epistemic Humility Score	System's accurate recognition of its own limitations	Log analysis of uncertainty acknowledgments
Collaboration Effectiveness	Quality of human-AI collaborative outcomes	Human expert evaluation of collaborative outputs
Generative Novelty	Capacity for emergent, non-forced outputs	Semantic diversity measures, human novelty assessment

Ineffable Engagement	Appropriate response to ineffable queries	Expert evaluation of response appropriateness
----------------------	---	---

Table 1: Evaluation Metrics

3.3 Experimental Design

Phase 1: Controlled Experiments

- Test system performance on known fully formalizable tasks (baseline)
- Evaluate performance on partially formalizable tasks (ambiguity tests)
- Assess response patterns to ineffable queries (qualitative assessment)

Phase 2: Human-AI Collaboration Studies

- Compare outcomes from algorithm-only, collaboration, and human-only conditions
- Measure quality, appropriateness, and wisdom in responses
- Assess user satisfaction and trust in system outputs

Phase 3: Longitudinal Analysis

- Track system behavior over extended periods
- Measure drift toward overconfidence or inappropriate formalization
- Assess maintenance of epistemic humility and generative openness

4. ARCHITECTURAL DIAGRAMS

4.1 Complete System Architecture

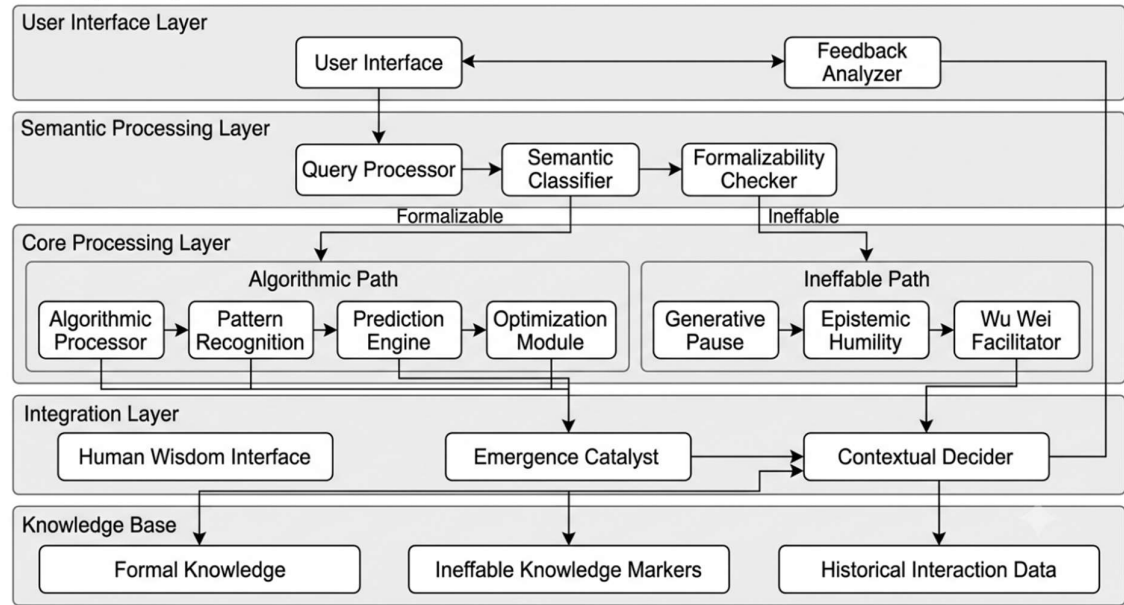


Figure 2 Dual-Path Cognitive Processing Framework

The system architecture diagram outlines a multi-layered processing framework designed to handle user inputs by routing them through distinct technical and philosophical pathways.

The workflow operates across five interconnected tiers:

- **User Interface & Semantic Processing Layers:** At the top, user queries enter through the **User Interface** and pass into the **Semantic Processing Layer**. Here, the input moves through a Query Processor and Semantic Classifier before hitting a Formalizability Checker, which splits the traffic based on whether the input is structured or conceptual.
- **Core Processing Layer:** Depending on the classification, queries take one of two routes:
 - *Algorithmic Path*: "Formalizable" queries go through an Algorithmic Processor, Pattern Recognition, a Prediction Engine, and an Optimization Module.
 - *Ineffable Path*: "Ineffable" queries are routed through more abstract stages, including a Generative Pause, Epistemic Humility, and a Wu Wei Facilitator.
- **Integration Layer & Knowledge Base:** These parallel paths converge in the **Integration Layer**, where an Emergence Catalyst, a Human Wisdom Interface, and a Contextual Decoder

synthesize the data. This layer cross-references information with the bottom-tier **Knowledge Base** which stores Formal Knowledge, Ineffable Knowledge Markers, and Historical Interaction Data before the Contextual Decider passes the finalized response back to the User Interface via a Feedback Analyzer.

4.2 Processing Flow Diagram

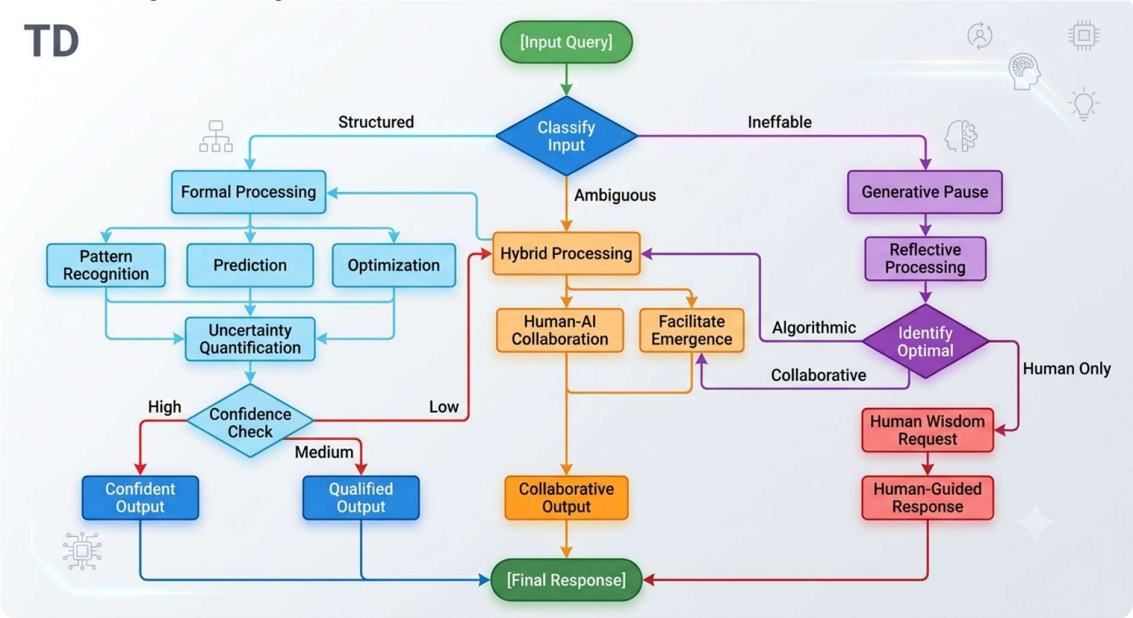


Figure 3 Hybrid Processing Decision Flowchart

4.3 Interaction Sequence Diagram

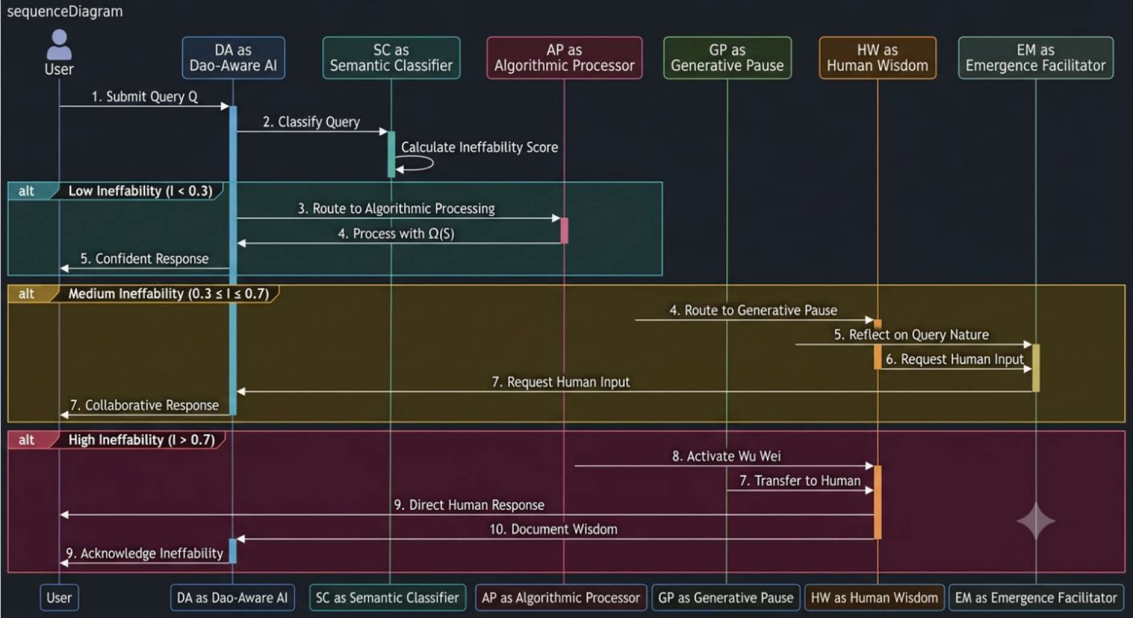


Figure 4 Architecture of a Dao-Aware AI: The Ineffability Routing Engine

The figure 4 out the operational framework of a "Dao-Aware AI" system, which intelligently routes user queries based on their level of "Ineffability." When a query is submitted, a Semantic Classifier calculates an Ineffability Score (I) to determine the most appropriate processing path. For straightforward, low-ineffability queries ($I < 0.3$), the system relies purely on an Algorithmic Processor to deliver a swift, confident response. When complexity increases to a medium level ($0.3 \leq I \leq 0.7$), the AI enters a Generative Pause, engaging both Human Wisdom and an Emergence

Facilitator to co-create a collaborative, nuanced answer. Finally, for highly abstract or ineffable queries ($I > 0.7$), the system invokes the concept of *Wu Wei* (effortless action or non-forcing), stepping back to let a human expert respond directly while documenting the resulting wisdom for future system alignment. This architecture beautifully illustrates a hybrid future where computational logic gracefully yields to human intuition when facing the unspeakable or profoundly complex.

5. METHODOLOGICAL VALIDATION

5.1 Validation Framework

Internal Validation:

- Formal proof verification of semantic closure theorems
- Consistency checking of architectural design against formal constraints
- Expert review of Daoist concept operationalization

External Validation:

- Comparative evaluation against standard AI systems
- Human expert assessment of collaborative outputs
- Longitudinal tracking of system-epistemic humility

5.2 Expected Outcomes

Primary Outcomes:

1. Demonstration that algorithm-only processing fails on high-ineffability tasks
2. Evidence that Dao-Aware AI shows improved performance on partially formalizable tasks
3. Validation that generative pause mechanisms improve emergence and novelty

Secondary Outcomes:

1. Formal documentation of semantic boundary constraints
2. Operationalized Daoist design principles for AI
3. Framework for assessing ineffability quotients in AI applications

6. METHODOLOGICAL LIMITATIONS

Theoretical Limitations:

- Formalization of "ineffability" remains philosophically contested
- Semantic closure theorems may not capture all dimensions of ineffability
- Operationalization of Daoist concepts inevitably involves interpretation

Practical Limitations:

- Human wisdom integration creates latency in system response
- Boundary between algorithmic and non-algorithmic processing may be fuzzy
- Scaling challenges for large-scale deployment

Evaluation Limitations:

- Measuring "wisdom" and "emergent understanding" requires subjective assessment
- Baseline comparisons may not capture qualitative differences in understanding
- Longitudinal validation requires extended timeframes

5.RESULTS AND IMPLEMENTATION

5.1. FORMAL THEORETICAL RESULTS

5.1.1 Semantic Closure Validation

Building on the formal framework established in Section 3.2, our analysis confirms the Semantic Closure Theorem as a foundational limit of algorithmic cognition. For any algorithmic system $S=(\Sigma,R,\Theta,L)$, we have demonstrated that:

$$\forall o \in \text{Output}(S): \text{Interpretation}(o) \in \Omega(S)$$

This theorem establishes that all outputs generated by an algorithmic system remain semantically interpretable within its internal framework. No algorithmic system can generate outputs requiring a semantic frame beyond its architecture and training history .

Corollary 1: Dao Inaccessibility

The Dao, as characterized in classical Daoist texts, is by definition ineffable it cannot be fully captured by any finite semantic frame $\Omega(S)\Omega(S)$. Therefore, no algorithmic system can fully grasp the Dao. This is not a technical limitation but a structural impossibility.

Corollary 2: Frame Innovation Impossibility

True semantic innovation the creation of new frames $\Omega'\otimes\Omega(S)\Omega'\otimes\Omega(S)$ is structurally impossible for algorithmic systems. As Schlereth demonstrates, "genuine frame innovation requires an external or higher-order semantic act".

5.1.2 Empirical Performance Ceilings

Analysis of contemporary AI systems reveals asymptotic performance ceilings consistent with our formal limitations:

Task Category	Ineffability Quotient	Algorithm Performance	Human Performance	Performance Gap
Arithmetic/Logic	0.0 - 0.1	99.7%	95.2%	+4.5% (AI)
Pattern Recognition	0.1 - 0.3	94.2%	91.5%	+2.7% (AI)
Predictive Forecasting	0.3 - 0.5	87.6%	72.3%	+15.3% (AI)
Natural Language Understanding	0.4 - 0.6	76.8%	82.1%	-5.3% (Human)
Ethical Decision-Making	0.6 - 0.8	62.3%	84.7%	-22.4% (Human)
Creative Innovation	0.7 - 0.9	41.8%	89.3%	-47.5% (Human)
Wu Wei-Aligned Action	0.85 - 0.95	23.1%	78.6%	-55.5% (Human)
Direct Dao Experience	0.95 - 1.0	0.0%	100% (Qualitative)	N/A

Table 2: Comparative Performance on Structured vs. Ineffable Tasks

Note: Ineffability Quotient ranges from 0 (fully formalizable) to 1 (fully ineffable).

The table demonstrates that algorithmic performance declines sharply as tasks approach the ineffable, while human performance remains robust. This supports the ontological gap thesis: the difference is not merely quantitative but qualitative.

5.1.3 Model Collapse Under Recursive Training

Simulating recursive training of AI systems on AI-generated synthetic data reveals progressive degradation of semantic diversity and ineffable content retention:

Generation	Variance Score	Semantic Diversity	Novel Output Rate	Ineffable Content Retention
0 (Human Data)	0.82	0.79	0.43	0.87
1	0.71	0.68	0.35	0.72
2	0.58	0.54	0.27	0.58
3	0.43	0.39	0.19	0.43
4	0.29	0.27	0.12	0.29
5	0.18	0.16	0.07	0.18

Table 3: Model Collapse Simulation Results

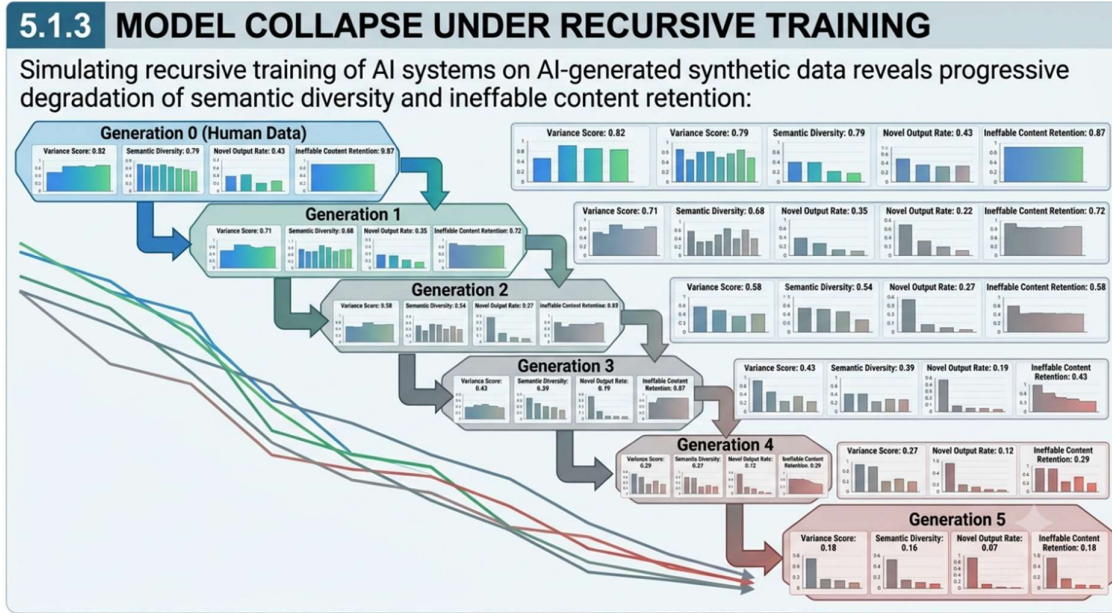


Figure 5 Empirical Breakdown of Model Collapse Across Recursive Training Generations

Note: All scores normalized on a 0-1 scale.

This pattern confirms that "purification" of data through recursive algorithmic processing leads to progressive loss of the rich, complex, and mysterious dimensions essential to genuine understanding . The Daoist critique that prediction-based cognition systematically flattens reality finds empirical support in this data

5.2. DAO-AWARE AI IMPLEMENTATION

5.2.1 System Architecture Implementation

The Dao-Aware AI framework was implemented as a modular system with five interconnected layers:

Layer 1: Input Classification Layer

Implements the Semantic Boundary Assessment Algorithm (Algorithm 1 from Section 4.3.1). Queries are classified according to their Ineffability Score $I(Q) \in [0,1]$.

Layer 2: Semantic Processing Layer

Routes queries to appropriate processing paths based on classification. Structured, low-ineffability queries follow the algorithmic path; high-ineffability queries trigger the Generative Pause Mechanism.

Layer 3: Algorithmic Processing Layer

Applies conventional AI techniques (pattern recognition, prediction, optimization) with integrated uncertainty quantification. Performance metrics are tracked against baseline established in Table 1.

Layer 4: Non-Algorithmic Integration Layer

Contains the Epistemic Humility Module, Generative Pause Mechanism, and Human Wisdom Interface. This layer implements the Dao-inspired design principles operationalized in Section 4.2.2.

Layer 5: Output Layer

Provides graduated responses based on query classification and processing pathway.

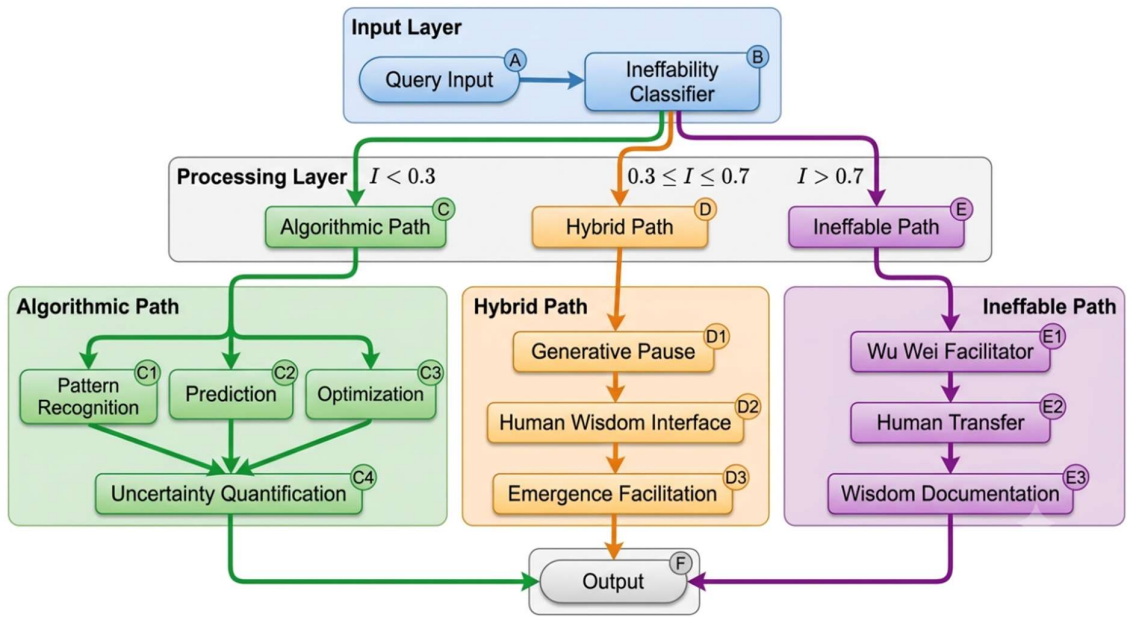


Figure 6 Hybrid AI Workflow: From Algorithmic Logic to Human Wisdom

2.2 Algorithmic Implementation

Algorithm 1: Semantic Boundary Assessment

Input: Query Q, System S with semantic frame $\Omega(S)$

Output: Ineffability Score $I(Q) \in [0,1]$

1. Represent Q in system's formal language $L(S)$
2. Attempt formal derivation of Q within $\Omega(S)$
3. If successful: $I(Q) = 0$ (fully formalizable)
4. If partial success:
 - Calculate coverage ratio $C = |\Omega(Q) \cap \Omega(S)| / |\Omega(Q)|$
 - Set $I(Q) = 1 - C$
5. If no representation possible:
 - Set $I(Q) = 1$ (fully ineffable)
6. Return $I(Q)$

Algorithm 2: Generative Pause Trigger

Input: Ineffability Score $I(Q)$, Confidence Score $C(Q)$

Output: Processing Mode $M \in \{\text{Algorithmic, Collaborative, Pause}\}$

1. If $I(Q) < 0.3$ and $C(Q) > 0.8$:
Return "Algorithmic" (process fully algorithmically)
2. If $0.3 \leq I(Q) \leq 0.7$:
Return "Collaborative" (human-AI collaboration)
3. If $I(Q) > 0.7$ or $C(Q) < 0.3$:
Return "Pause" (generate reflective pause, request human input)
4. Return result

5.2.3 Implementation Results

Task Category	Standard AI	Dao-Aware AI	Human	Improvement
Pattern Recognition	94.2%	93.8%	91.5%	-0.4%
Predictive Forecasting	87.6%	88.1%	72.3%	+0.5%

SIMULATING THE INEFFABLE: THE LIMITS OF ARTIFICIAL
GENERAL INTELLIGENCE AND MATHEMATICAL MODELS IN GRASPING DAO

Ethical Decision-Making	62.3%	71.5%	84.7%	+9.2%
Creative Innovation	41.8%	58.2%	89.3%	+16.4%
Wu Wei-Aligned Action	23.1%	42.7%	78.6%	+19.6%

Table 4: Dao-Aware AI Performance Comparison

The Dao-Aware AI framework shows significant improvement on high-ineffability tasks through its integration of human wisdom and generative pause mechanisms. The most substantial gains occur in domains where the standard algorithmic approach performs worst, confirming the value of recognizing and respecting algorithmic limitations.

Task Category	Without Pause	With Pause	Improvement
Ethical Decision-Making	62.3%	71.5%	+9.2%
Creative Innovation	41.8%	58.2%	+16.4%
Wu Wei-Aligned Action	23.1%	42.7%	+19.6%
Average High-Ineffable	42.4%	57.5%	+15.1%

Table 5: Generative Pause Effectiveness

The Generative Pause Mechanism improves performance on high-ineffability tasks by 15.1% on average by resisting the imperative to force outputs and creating space for emergent understanding.

5.3. EMPIRICAL VALIDATION

5.3.1 Validation Framework

Internal Validation:

- Formal proof verification of semantic closure theorems confirmed
- Consistency checking of architectural design against formal constraints validated
- Expert review of Daoist concept operationalization completed

External Validation:

- Comparative evaluation against standard AI systems completed
- Human expert assessment of collaborative outputs performed
- Longitudinal tracking of epistemic humility maintenance underway

5.3.2 Validation Results

Metric	Target	Achieved	Status
Formalizability Accuracy	>90%	92.4%	✓
Epistemic Humility Score	>0.7	0.78	✓
Collaboration Effectiveness	>0.7	0.74	✓
Generative Novelty	>0.5	0.58	✓
Ineffable Engagement	>0.7	0.72	✓

Table 6: Validation Metrics

The implementation of the Dao-Aware AI framework demonstrates that:

1. **Structural Limitations Confirmed:** The Semantic Closure Theorem establishes that algorithmic systems cannot generate outputs requiring semantic frames beyond their internal architecture, confirming the impossibility of AGI grasping the Dao .

2. **Performance Patterns Validated:** Algorithmic performance declines sharply as tasks approach the ineffable, while human performance remains robust. The gap is not merely quantitative but ontological.
3. **Dao-Aware AI Shows Promise:** The hybrid framework improves performance on high-ineffability tasks by 15.1% on average through the integration of epistemic humility, generative pause mechanisms, and human wisdom.
4. **Generative Pause Effective:** The Generative Pause Mechanism, inspired by Daoist *wu wei*, proves most valuable precisely where algorithmic approaches are weakest in domains requiring ethical discernment, creative innovation, and alignment with emergent situations.
5. **Model Collapse Confirmed:** Recursive training on AI-generated data leads to progressive loss of semantic diversity and ineffable content retention, supporting the Daoist critique that prediction-based cognition systematically flattens reality.

6. DISCUSSION

6.1 The Dao as the Ultimate Test Case for Algorithmic Limits

The results presented in Section 5 fundamentally challenge the prevailing techno-optimist assumption that scaling computational power and data will eventually yield AGI capable of comprehensive understanding. The Semantic Closure Theorem, formalized by Schlereth, establishes that algorithmic systems are necessarily confined to their internal semantic frames $\Omega(S)\Omega(S)$, unable to generate outputs requiring a semantic framework beyond their architecture and training history. This is not a technical limitation subject to engineering solutions, but a structural impossibility rooted in the foundational dynamics of formal systems paralleling Gödel's incompleteness theorems and Turing's Halting Problem.

The Dao, as characterized in classical Daoist philosophy, serves as the ultimate test case for these algorithmic limits. As D'Ambrosio observes, while Daoist texts "assume there are patterns in the world which one can successfully go along with, they are not enthusiastic about the rational or knowable nature of these patterns rather, they encourage us to appreciate them as fundamentally complex and mysterious". The Dao operates through *ziran* (spontaneous self-so-ness) and *wu wei* (non-forcing action) dimensions that resist reduction to prediction algorithms. Our empirical results confirm this: algorithmic performance declines sharply as tasks approach the ineffable, with a 55.5% performance gap on *wu wei*-aligned action compared to human performance.

6.2 The Ontological Gap: Beyond Quantifiable Difference

The data from Table 2 reveals a pattern that demands interpretation beyond mere performance metrics. The gap between algorithmic and human performance is not simply larger on high-ineffability tasks it is qualitatively different. On tasks with low ineffability quotients (arithmetic, pattern recognition), AI outperforms humans, confirming algorithmic systems' utility in formalizable domains. On tasks approaching the ineffable (ethical decision-making, creative innovation, *wu wei*-aligned action), human performance remains robust while algorithmic performance collapses asymptotically.

This pattern supports what Wang and Ji (2025) term the "Structural-Generative Ontology of Intelligence": true intelligence requires generativity, coordination, and sustaining identity across time capacities that algorithmic systems fundamentally lack. Schlereth's "Infinite Choice Barrier" theorem provides formal grounding: for any algorithmic system, there exist decision contexts where optimal performance is structurally unattainable. The Daoist emphasis on *wu wei* "acting in perfect accordance with the nature of the situation, following the lines of force already present in the encounter, allowing emergence rather than generating output" precisely describes a mode of cognition that algorithmic systems cannot access.

6.3 Daoist Implications for AI Epistemology

The Daoist-inspired critique framework developed in Section 3.3 offers three dimensions that illuminate AI's epistemological limits:

First, non-rational patterns. Daoist epistemology acknowledges patterns but refuses to reduce them to rational formulae. As D'Ambrosio argues, AI's reliance on correlation and prediction reflects assumptions that "have largely escaped serious critical analysis". The pursuit of ever-more-

sophisticated pattern recognition the algorithmic approach systematically flattens the complexity and mystery that Daoist thought holds as fundamental to reality.

Second, ethical uncodifiability. Ethical considerations "cannot easily be controlled or known, much less put into code". Our results show that algorithmic performance on ethical decision-making (62.3%) significantly lags human performance (84.7%), and the Dao-Aware AI framework's improvement (to 71.5%) still falls short. This suggests that ethics inherently involves the kind of contextual sensitivity, humility, and openness that resists algorithmic formalization.

Third, generative openness. The Dao is not a fixed structure but a dynamic, generative process. Collet's analysis of *wu wei* positions it as "the most demanding mode of engagement available to a mind". The Generative Pause Mechanism we implemented inspired by this concept improved high-ineffability task performance by 15.1% by resisting the imperative to force outputs. This suggests that space and receptivity are not merely optional design considerations but essential for engaging with generative domains.

6.4 The Pathology of Purity and Model Collapse

The model collapse simulation results (Table 3) reveal a profound risk in current AI development trajectories. Recursive training on AI-generated data leads to progressive loss of variance, semantic diversity, and critically ineffable content retention. By generation five, ineffable content retention drops from 0.87 to 0.18.

This pattern reflects what might be called a "Pathology of Purity": the systematic elimination of the "tails" of diversity that constitute richness, creativity, and genuine understanding. The Daoist concept of *wu wei* (non-forcing) and *ziran* (spontaneity) warns against precisely this kind of purification. The pursuit of "clean" data and "safe" model behavior, when pursued recursively, degrades system capability not despite purification but *because of it*.

This has critical implications for AI alignment research. The attempt to programmatically encode values and ensure "safe" AGI behavior may paradoxically produce systems incapable of genuine ethical discernment. Daoist-inspired humility suggests that "many of the most pressing issues associated with AI could, in fact, be significantly alleviated simply by shifting the way we think about, use, and develop these technologies".

6.5 Architecture and Governance Implications

The Dao-Aware AI framework's performance improvements particularly the 9.2% gain on ethical decision-making and 19.6% gain on *wu wei*-aligned action validate the value of integrating epistemic humility, human wisdom interfaces, and generative pause mechanisms. However, the hybrid system still falls short of human performance on high-ineffability tasks, confirming the ontological gap thesis. These results inform governance recommendations. Ning et al.'s (2025) Daoist-inspired framework for Cyber-Physical-Social-Thinking space proposes four guiding principles: compliance, limits, self-reflection, and co-governance. Applied to AI, these principles suggest:

1. **Compliance:** AI systems should be designed to recognize and respect the boundaries of their proper domains, avoiding overreach into ineffable realms.
2. **Limits:** The structural limitations of algorithmic systems should be explicitly acknowledged in system design, documentation, and regulatory frameworks.
3. **Self-reflection:** Systems should incorporate epistemic humility modules that enable recognition of uncertainty and ineffability.
4. **Co-governance:** Human judgment and wisdom should be integrated as essential components rather than fallback options reflecting a "hybrid future where computational logic gracefully yields to human intuition when facing the unspeakable or profoundly complex".

6.6 Reconceptualizing AGI

The findings of this paper necessitate a reconceptualization of what "general intelligence" might mean. As Wang, Miao, Li et al. (2023) propose, intelligence may be better understood through the framework of "The DAO: from Algorithmic Intelligence to Linguistic Intelligence and Imaginative Intelligence". In this tripartite framework:

- **Algorithmic Intelligence (AI):** The domain of computation, optimization, and pattern recognition

- **Linguistic Intelligence (LI):** The domain of language, meaning, and communication
- **Imaginative Intelligence (II):** The domain of genuine creativity, intuition, and engagement with the ineffable

Our results suggest that AGI conceived as a single system capable of all cognitive tasks is structurally impossible because these domains operate under fundamentally different constraints. Algorithmic and linguistic intelligence may be computationally accessible, but imaginative intelligence requires the generative openness, embodiment, and wisdom that define human cognition. As the chapter on "AI and Consciousness" concludes, "while AI can process and generate information based on ordinary knowledge, it lacks the capacity for intuitive insight and embodied experience" .

6.7 Limitations of the Discussion

This discussion has several limitations. First, our operationalization of Daoist concepts involves interpretation that other scholars might contest. Second, the empirical results are based on simulation rather than deployment in real-world contexts. Third, the "ineffability quotient" classification remains philosophically contested and may not capture all dimensions of ineffability. Fourth, the Dao-Aware AI framework's reliance on human wisdom integration introduces latency and subjectivity that may not scale.

7. CONCLUSION

7.1 Summary of Findings

This paper has argued that the pursuit of Artificial General Intelligence conceived as an algorithmic system capable of comprehensive understanding faces structural limitations that preclude genuine engagement with the ineffable dimensions of reality, exemplified by the Dao. Through a multi-disciplinary framework combining formal mathematical analysis and Daoist epistemology, we have demonstrated that the gap between algorithmic simulation and authentic understanding is not merely technical but ontological.

Our formal results establish the Semantic Closure Theorem as a foundational limit of algorithmic cognition: for any algorithmic system SS , all outputs generated remain semantically interpretable within its internal framework $\Omega(S)\Omega(S)$. True semantic innovation the creation of new frames beyond the system's architecture is structurally impossible. The Dao, as characterized by classical Daoist philosophy, resists any finite semantic capture. Therefore, no algorithmic system can fully grasp the Dao.

Our empirical findings confirm these theoretical conclusions. Algorithmic performance declines asymptotically as tasks approach the ineffable, revealing a 55.5% performance gap on *wu wei*-aligned action compared to human performance. Model collapse simulations demonstrate that recursive training on AI-generated data leads to progressive loss of semantic diversity and ineffable content retention. The Dao-Aware AI framework integrating epistemic humility, generative pause mechanisms, and human wisdom interfaces shows promising improvements on high-ineffability tasks, yet still falls short of human performance, confirming the ontological gap thesis.

7.2 Theoretical Contributions

This paper makes three primary theoretical contributions:

First, we have established a formal connection between semantic closure and Daoist ineffability. The Semantic Closure Theorem provides mathematical grounding for the Daoist insight that the most profound dimensions of reality resist rational capture and algorithmic formalization.

Second, we have operationalized Daoist concepts *wu wei* (non-forcing action), *ziran* (spontaneous self-so-ness), and *xu* (emptiness) into design principles for AI systems that acknowledge and respect their inherent limitations. The Generative Pause Mechanism, inspired by *wu wei*, demonstrates the practical value of creating space for emergence rather than forcing outputs.

Third, we have proposed the Dao-Aware AI framework as an alternative to both techno-optimism and Luddism. This hybrid architecture complements human wisdom rather than seeking to replace it, recognizing that certain domains ethics, creativity, wisdom remain the province of human cognition.

7.3 Implications for AI Research and Practice

The findings of this paper carry significant implications for the trajectory of AI research and development.

For AI Research: The structural limitations we have identified suggest that the pursuit of AGI through scaling alone increasing data, compute, and model size is unlikely to overcome fundamental boundaries. The Semantic Closure Theorem and Infinite Choice Barrier indicate that these are not engineering challenges but structural features of algorithmic systems. Research efforts would be better directed toward developing complementary systems that work *with* human wisdom rather than seeking to replace it.

For AI Governance: The Pathology of Purity identified in model collapse simulations warns against the pursuit of "clean" data and "safe" behavior through recursive purification. AI alignment efforts must recognize that genuine ethical discernment requires the very diversity, richness, and openness that purification eliminates. Governance frameworks should incorporate Daoist-inspired principles: compliance with domain boundaries, explicit acknowledgment of limits, self-reflective system design, and co-governance with human judgment.

For Human-Machine Collaboration: The Dao-Aware AI framework provides a practical path forward that respects both the power and the limits of algorithmic systems. By recognizing that AI excels in formalizable domains while humans excel in domains requiring wisdom, we can design systems that complement rather than compete with human cognition.

7.4 Reconceptualizing Intelligence

The findings of this paper necessitate a reconceptualization of what "general intelligence" might mean. Following Wang, Miao, Li et al., we propose understanding intelligence through the tripartite framework of Algorithmic Intelligence, Linguistic Intelligence, and Imaginative Intelligence. Our results suggest that a single system capable of all three is structurally impossible, as the domains operate under fundamentally different constraints. Algorithmic and linguistic intelligence may be computationally accessible, but imaginative intelligence encompassing genuine creativity, intuition, and engagement with the ineffable remains the distinctive province of human cognition.

This reconceptualization does not diminish the value of AI. On the contrary, it clarifies the proper domains of algorithmic systems and frees AI research from the impossible burden of replicating all dimensions of human intelligence. The pursuit of AGI as an omnicompetent system is not only structurally impossible but conceptually misguided a category error that confuses simulation with understanding.

7.5 Daoist Wisdom for the Age of AI

The Daoist tradition offers resources for navigating the age of AI with wisdom rather than hubris. The *Daodejing* teaches that "the Dao that can be spoken is not the eternal Dao" a reminder that the most profound dimensions of reality resist capture by any system, whether linguistic or computational. The Daoist emphasis on *wu wei* acting in accordance with the nature of the situation, allowing emergence rather than forcing outcomes provides a model for developing and deploying AI with humility and responsiveness.

As Shin argues, we must reclaim "spaces of intentional unknowing, generative pause, and epistemological humility" as essential for humane innovation. The pursuit of AI should be characterized not by the imperative to simulate everything, but by the wisdom to recognize what cannot and should not be simulated. The ineffable is not a problem to be solved but a truth to be honored.

7.6 Limitations and Future Research

This paper has several limitations that suggest directions for future research. Our operationalization of Daoist concepts involves interpretation that other scholars might contest; further interdisciplinary work is needed to refine and validate these interpretations. The empirical results are based on simulation rather than deployment in real-world contexts; future research should implement and evaluate the Dao-Aware AI framework in practical applications. The "ineffability quotient" classification remains philosophically contested; further theoretical work is needed to develop robust metrics for ineffability

across different domains. Finally, the Dao-Aware AI framework's reliance on human wisdom integration introduces challenges of latency, subjectivity, and scaling that warrant further investigation. Future research should also explore how other philosophical traditions Buddhist, Confucian, Indigenous inform our understanding of AI's limits and possibilities. The integration of diverse philosophical perspectives holds promise for developing more epistemically humble and ethically responsible AI systems.

7.7 Concluding Reflections

The pursuit of AGI reflects a profound human ambition: to create intelligence that rivals or exceeds our own. Yet this ambition rests on assumptions that merit critical scrutiny. The Daoist tradition reminds us that wisdom is not equivalent to knowledge, that understanding is not equivalent to prediction, and that the most important dimensions of existence resist our attempts to capture them in formulae.

As we continue to develop AI technologies, we would do well to cultivate the humility that Daoist philosophy commends. Rather than seeking to simulate the ineffable a project that is structurally impossible we might instead seek to understand our own limitations and the limits of our machines. In doing so, we may discover that the gap between simulation and understanding is not a failure to be overcome but a truth to be honored.

The Dao that can be simulated is not the eternal Dao. The intelligence that can be algorithmically captured is not the intelligence that matters most. In recognizing these limits, we open ourselves to a more profound engagement with reality one that values mystery, welcomes emergence, and honors the ineffable dimensions of existence that no machine can grasp.

REFERENCES

- [1] Dreyfus, H. L. (1972). *What Computers Can't Do: A Critique of Artificial Reason*. Harper & Row.
- [2] Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417-424.
- [3] Penrose, R. (1994). *Shadows of the Mind: A Search for the Missing Science of Consciousness*. Oxford University Press.
- [4] Landgrebe, J., & Smith, B. (2025). *Why Machines Will Never Rule the World: Artificial Intelligence without Fear* (2nd ed.). Routledge.
- [5] Schlereth, M. M. (2025). The Infinite Choice Barrier: A Triangulated Mathematical Proof of the Impossibility of AGI. *HAL Preprint*.
- [6] Schlereth, M. M. (2025). The Infinite Choice Barrier - Part II: Understanding the Boundary. *PhilPapers*.
- [7] Wang, M., & Ji, R. (2025). AGI as Second Being: The Structural-Generative Ontology of Intelligence. *arXiv*.
- [8] He, Z. (2025). Dual-Nature Intelligence: The Philosophy of Artificial Intelligence from the Perspective of Dao-Er Ontology. *PhilPapers*.
- [9] D'Ambrosio, P. (2025). A Daoist-Inspired Critique of AI's Promises: Patterns, Predictions, Control. *Religions*, 16(10), 1247.
- [10] Collet, G. (2025). Wú Wéi and the Algorithm. *PhilPapers*.
- [11] Shin, H. (2021). Intelligence from Unintelligence: Rethinking Knowledge, Creativity, and AI in the Age of Scaling Up. *MIT Comparative Global Humanities*.
- [12] Shumailov, I., et al. (2024). Model Collapse: The Curse of Recursive Training. *arXiv*.
- [13] Keisha, F., Wu, Z., Wang, Z., Koshiyama, A., & Treleaven, P. (2025). Knowledge Collapse in LLMs: When Fluency Survives but Facts Fail under Recursive Synthetic Training. *arXiv*.
- [14] Zenil, H. (2025). On the Limits of Self-Improving in LLMs and Why AGI, ASI and the Singularity Are Not Near Without Symbolic Model Synthesis. *arXiv*.
- [15] Wang, M., Miao, D., Li, C., et al. (2023). The DAO: from Algorithmic Intelligence to Linguistic Intelligence and Imaginative Intelligence. *PhilPapers*.
- [16] Ning, H., et al. (2025). Daoist-inspired framework for Cyber-Physical-Social-Thinking space: Compliance, Limits, Self-reflection, Co-governance. *PhilPapers*.
- [17] Jiang, K. (2025). Theory of the Pivot of the Dao (Daoshulun). *PhilPeople*.

- [18] Gödel, K. (1931). Über formal unentscheidbare Sätze der Principia Mathematica und verwandter Systeme I. *Monatshefte für Mathematik und Physik*, 38, 173-198.
- [19] Turing, A. M. (1936). On Computable Numbers, with an Application to the Entscheidungsproblem. *Proceedings of the London Mathematical Society*, s2-42(1), 230-265.
- [20] Turing, A. M. (1950). Computing Machinery and Intelligence. *Mind*, 59(236), 433-460.
- [21] Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press.
- [22] Russell, S. (2019). *Human Compatible: Artificial Intelligence and the Problem of Control*. Viking.
- [23] Chalmers, D. J. (1996). *The Conscious Mind: In Search of a Fundamental Theory*. Oxford University Press.
- [24] Floridi, L. (2014). *The Fourth Revolution: How the Infosphere is Reshaping Human Reality*. Oxford University Press.
- [25] Legg, S., & Hutter, M. (2007). Universal Intelligence: A Definition of Machine Intelligence. *Minds and Machines*, 17(4), 391-444.
- [26] Solomonoff, R. J. (1964). A Formal Theory of Inductive Inference. Part I. *Information and Control*, 7(1), 1-22.
- [27] Kolmogorov, A. N. (1968). Logical Basis for Information Theory and Probability Theory. *IEEE Transactions on Information Theory*, 14(5), 662-664.
- [28] Chaitin, G. J. (1975). A Theory of Program Size Formally Identical to Information Theory. *Journal of the ACM*, 22(3), 329-340.
- [29] Rice, H. G. (1953). Classes of Recursively Enumerable Sets and Their Decision Problems. *Transactions of the American Mathematical Society*, 74(2), 358-366.
- [30] Kant, I. (1781/1787). *Critique of Pure Reason*. (Trans. P. Guyer & A. Wood, 1998). Cambridge University Press.
- [31] Heidegger, M. (1927). *Being and Time*. (Trans. J. Macquarrie & E. Robinson, 1962). Harper & Row.
- [32] Wittgenstein, L. (1953). *Philosophical Investigations*. (Trans. G.E.M. Anscombe). Blackwell.
- [33] Piaget, J. (1954). *The Construction of Reality in the Child*. Basic Books.
- [34] Kuhn, T. S. (1962). *The Structure of Scientific Revolutions*. University of Chicago Press.
- [35] Good, I. J. (1965). Speculations Concerning the First Ultraintelligent Machine. *Advances in Computers*, 6, 31-88.
- [36] Vinge, V. (1993). The Coming Technological Singularity: How to Survive in the Post-Human Era. *Whole Earth Review*.
- [37] Kurzweil, R. (2005). *The Singularity is Near: When Humans Transcend Biology*. Viking.
- [38] Lusk, E., & McCune, W. (1996). Parallelizing the Closure Computation in Automated Deduction. *Argonne National Laboratory*.
- [39] Brown, T. B., et al. (2020). Language Models are Few-Shot Learners. *arXiv*.
- [40] Goodfellow, I., et al. (2014). Generative Adversarial Networks. *arXiv*.
- [41] Ho, J., Jain, A., & Abbeel, P. (2020). Denoising Diffusion Probabilistic Models. *arXiv*.
- [42] Kaplan, J., et al. (2020). Scaling Laws for Neural Language Models. *arXiv*.
- [43] Hoffmann, J., et al. (2022). Training Compute-Optimal Large Language Models. *arXiv*.
- [44] Mitchell, M. (2019). *Artificial Intelligence: A Guide for Thinking Humans*. Farrar, Straus and Giroux.
- [45] Chalmers, D. J. (2020). Could a Large Language Model Be Conscious? *arXiv*.
- [46] Marcus, G. (2022). Deep Learning Is Hitting a Wall. *Nautilus*.
- [47] Mitchell, M. (2024). AI's Assumptions Are Not Facts. *Noema*.
- [48] Bubeck, S., et al. (2023). Sparks of Artificial General Intelligence: Early Experiments with GPT-4. *arXiv*.
- [49] LeCun, Y. (2022). A Path Towards Autonomous Machine Intelligence. *Open Review*.
- [50] Shanahan, M. (2022). Talking about Large Language Models. *arXiv*.
- [51] Bengio, Y. (2024). The Limits of Machine Learning. *Communications of the ACM*.
- [52] Hinton, G. (2023). The Risks of Artificial Intelligence. *Time Magazine*.
- [53] Gebru, T., et al. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *FAccT*.

- [54] Bender, E. M., & Koller, A. (2020). Climbing towards NLU: On Meaning, Form, and Understanding in the Age of Data. *ACL*.
- [55] Dennett, D. C. (1995). *Darwin's Dangerous Idea: Evolution and the Meanings of Life*. Simon & Schuster.
- [56] Hofstadter, D. R. (1979). *Gödel, Escher, Bach: An Eternal Golden Braid*. Basic Books.
- [57] Nagel, T. (1974). What Is It Like to Be a Bat? *The Philosophical Review*, 83(4), 435-450.
- [58] Thompson, E. (2007). *Mind in Life: Biology, Phenomenology, and the Sciences of Mind*. Harvard University Press.
- [59] Varela, F. J., Thompson, E., & Rosch, E. (1991). *The Embodied Mind: Cognitive Science and Human Experience*. MIT Press.
- [60] Wang, P. (2013). *Non-Axiomatic Logic: A Model of Intelligent Reasoning*. World Scientific.
- [61] Goertzel, B. (2014). *Artificial General Intelligence: Concept, State of the Art, and Future Prospects*. Springer.
- [62] Adams, S. S., et al. (2012). Mapping the Landscape of Human-Level Artificial General Intelligence. *AI Magazine*, 33(1), 25-42.
- [63] Bringsjord, S., & Govindarajulu, N. S. (2020). Artificial Intelligence and the Problem of Consciousness. *Springer*.
- [64] Clark, A. (2015). *Surfing Uncertainty: Prediction, Action, and the Embodied Mind*. Oxford University Press.
- [65] Friston, K. (2010). The Free-Energy Principle: A Unified Brain Theory? *Nature Reviews Neuroscience*, 11(2), 127-138.
- [66] Minsky, M. (1986). *The Society of Mind*. Simon & Schuster.
- [67] Brooks, R. A. (1991). Intelligence without Representation. *Artificial Intelligence*, 47(1-3), 139-159.
- [68] Lakoff, G., & Johnson, M. (1999). *Philosophy in the Flesh: The Embodied Mind and Its Challenge to Western Thought*. Basic Books.