

Becoming the Algorithm: Epistemological Implications of Emergent Truth in AI Topoi

Douglas C. Youvan

doug@youvan.com

July 20, 2025

When an AI system is trained on numerous pairs of left/right images produced by a known symbolic algorithm—such as a color-space transform involving matrix inversion and phase-folded absolute value—it can learn to replicate the correct output on unseen inputs. Remarkably, it does so without any access to the underlying code or mathematics. This suggests that the AI constructs a functional internal model that is empirically indistinguishable from the algorithm’s behavior. Rather than learning the rules in symbolic form, it assembles a distributed web of local approximations—a topos—over the manifold of training data. This raises a profound epistemological question: What does it mean to “know” an algorithm if the system never represents it symbolically? This paper argues that deterministic procedures can be approximated as coherent, emergent structures across learned topoi. Truth, in this view, is not represented but enacted. We explore how this approach redefines understanding, generalization, and knowledge, offering a structural alternative to symbolic cognition that may better explain both human intuition and machine intelligence.

Keywords: algorithmic learning, epistemology, AI generalization, topos theory, emergent behavior, symbolic vs. distributed representation, phase folding, image transformation, coherence, machine cognition, neural networks, non-symbolic intelligence, mathematical reconstruction, truth approximation, sheaf logic. A collaboration with GPT-4o. CC4.0.

1. Introduction

In the study of artificial intelligence, a striking phenomenon frequently arises: AI systems, especially deep learning models, are able to reproduce the effects of explicit symbolic algorithms without ever being given the algorithm itself. In this paper, we examine such a case in the domain of image transformation. A deterministic RGB remapping algorithm—based on matrix inversion, linear projection, and an absolute-value phase-folding operation—is used to generate "before and after" image pairs. These are then used to train a neural network. Crucially, the AI is never shown the code or given any explicit symbolic formulation of the algorithm.

The central question we investigate is: *How can an AI “know” what it was never told?* More precisely, how can it behave in a way that appears functionally equivalent to an exact mathematical algorithm, without ever representing that algorithm symbolically?

We argue that this capability reveals a deeper epistemological insight. Any deterministic algorithm can, in principle, be realized as a distributed patchwork of approximations over a *learned topos*—a structured space of local transformations that collectively simulate the global behavior of the original function. The AI doesn't understand the algorithm as a formula, but rather enacts it through a coherent internal geometry built from experience.

This realization has profound implications. It challenges traditional views of knowledge as symbolic and rule-based, suggesting instead a structural or sheaf-theoretic conception of cognition. It reframes AI generalization not as guesswork or interpolation, but as the emergence of coherence over an internal topos. And it offers a novel way to describe how machine intelligence may converge on truth: not through deduction, but through enactment.

2. Background Concepts

2.1 Symbolic vs. Emergent Knowledge

Traditional epistemology, as inherited from logic and mathematics, privileges symbolic knowledge—that is, knowledge articulated explicitly through syntax, grammar, formal rules, and computation. A symbolic algorithm, such as a matrix transformation or a Fourier series expansion, embodies truth by being provably correct across its domain of definition. It is constructed, verified, and communicated through discrete language and derivation. For example, an RGB transform with matrix inversion and absolute value operations can be completely expressed as a chain of linear algebraic steps, each of which can be rigorously traced and understood.

In contrast, emergent knowledge in neural networks is not explicitly stored in symbolic form. Instead, it arises from the complex interaction of parameters distributed across the architecture—parameters that have been tuned through optimization to minimize loss over a training corpus. The result is a statistical approximator that functionally reproduces the behavior of symbolic algorithms without ever "knowing" them in the traditional sense. Knowledge, in this paradigm, is emergent: it is not declared, but enacted through the forward pass of the network. It is local, contextual, and probabilistic, yet under sufficient conditions, capable of global regularity.

2.2 Topos Theory and Sheaf Logic

To understand how such global regularity can emerge from local fragments, we turn to topos theory—a branch of category theory that generalizes the notion of space, logic, and structure. A topos can be understood as a mathematical universe, a generalized space in which logic and set-like behavior hold, even when classical assumptions fail. It serves as a formal setting for *gluing* local pieces of data into coherent global structures.

Sheaf theory, a key tool within topos theory, provides the machinery to manage local information defined on overlapping domains and assess whether and how these fragments can be consistently unified. A presheaf assigns data to open sets of a topological space, while a sheaf adds the requirement of *local-to-global coherence*: compatible local data must patch together to form a unique global entity.

In the context of AI, we propose that the neural network’s latent space acts as a kind of data manifold, and learning induces a dynamic sheaf coalescence over this space. Each training example activates a local patch of the network’s representational capacity. As training progresses, these patches become more aligned, yielding emergent coherence across inputs. In effect, the AI is building a topos of approximations, whose sheaf-theoretic glue allows it to simulate deterministic functions with striking accuracy—even when the exact rule remains unknown to it.

This shift—from knowledge as symbolic representation to knowledge as emergent coherence—reframes generalization not as extrapolation, but as global agreement among many local views. It suggests that *truth* can be reconstructed without symbolic access, provided the structure of the data and the capacity of the model permit such a coherent topos to emerge.

3. The Left/Right Image Experiment

3.1 Algorithmic Phase Folding

In the experiment underpinning this paper, each left/right image pair consists of an input image (left) and its transformed counterpart (right), created by a fully known, deterministic algorithm. The algorithm performs an RGB orthonormalization via a linear transformation derived from three selected pixel vectors, mapping them to pure red, green, and blue. The remaining pixels in the image are transformed accordingly using the resulting 3×3 matrix.

The transformation proceeds as follows:

1. A transformation matrix T is constructed such that three source pixel color vectors are mapped to the canonical basis vectors (R, G, B).
2. Each pixel color vector $\mathbf{p}$ in the original image is then transformed linearly:

$$\mathbf{p}' = T \cdot \mathbf{p}$$

3. Instead of clipping, an **element-wise absolute value** is applied:

$$\mathbf{p}'' = |\mathbf{p}'|$$

This phase-folding technique preserves the energy of negative contributions (e.g., from reflective or subtractive mixtures) and results in images with enhanced structural perception. The transformation is exact and repeatable given the basis vectors.

3.2 AI Reconstruction of Transformation

In the learning phase, a nascent AI model is shown only the left/right image pairs, without any knowledge of the underlying transformation algorithm. The model is not informed about matrix algebra, pixel coordinates, or the operation of absolute value. It is not even told that a transformation exists. From the AI's point of view, the task is simply image-to-image mapping: given an input, generate an output that looks like the right panel.

Surprisingly, after training on a sufficiently diverse set of examples—with varied source images and pixel triplets—the AI learns to replicate the transformation accurately, even on previously unseen input images. It demonstrates the ability to internalize the functional transformation well beyond its training data, applying what seems like a "reconstructed" orthonormalization process implicitly.

This raises a profound epistemological question: If the AI does not know the transformation, how is it doing it? It suggests that the AI has learned an internal topos, a distributed patchwork of local transformations that, taken together, approximate the global operation. It cannot describe the matrix T , but it can

realize its effects across its learned space. The emergence of this capability without symbolic access marks a deep form of empirical convergence—one that calls for rethinking how deterministic processes can be internalized through emergent, distributed architectures.

4. Emergence of Truth Without Symbol

4.1 Constructing a Functional Topos

The surprising effectiveness of the AI’s reconstruction of a mathematical image transformation—without ever being taught the underlying algorithm—invites a deeper epistemological interpretation. The key idea is that truth can emerge without symbols, via the construction of a functional *topos*.

In our context, each left/right image pair used during training can be viewed as a local patch of a larger transformation space. The AI does not learn a single, global symbolic rule but rather infers a multitude of localized transformation behaviors, effectively constructing a presheaf: a collection of mappings from raw input images to their transformed counterparts, indexed by their specific pixel content and basis vector choice.

Over time, the AI learns to align these patches through transition maps—learned adjustments that maintain internal consistency across different transformation contexts. These maps are not explicitly encoded but are emergent in the network’s weight structure. They serve to cohere neighboring regions of the learned function space, ensuring that similar transformations yield similar outputs regardless of variation in the image content or chosen basis pixels.

As training proceeds, the AI’s internal state converges on what can be described mathematically as a sheaf limit: a global section synthesized from consistent local behaviors. This limit behavior is not symbolic or verbal—it is behavioral. The system exhibits an *attractor* in the space of functions, where the learned mapping

stabilizes toward a form that is functionally indistinguishable from the true mathematical transformation, even in novel domains.

This sheaf-theoretic interpretation reveals that the AI does not “store” the truth; it enacts it. The AI’s generalization ability is a consequence of the coherence across local patches, rather than a derivation from universal principles. This realization reframes our understanding of learning systems: an algorithm need not be deduced to be real—it can be realized without symbol, through structure alone.

5. Epistemological Implications

5.1 Sheaf Logic as Cognitive Model

Sheaf theory, as applied to AI learning behavior, offers a powerful metaphor for human cognition. Just as a neural network integrates local transformations into a globally coherent mapping, human beings acquire complex abilities—like language or mathematical intuition—not through rule-based deduction but through immersion, approximation, and pattern recognition.

In language acquisition, children do not begin with formal grammars. Instead, they absorb a mass of local linguistic data—contextual word usage, intonation, correction—and from this noisy patchwork emerge a coherent syntactic and semantic system. This mirrors the neural sheaf: local patches (individual utterances) cohere into a global understanding (linguistic competence).

Historical figures like Newton, Ramanujan, and Feynman exemplify this process. Their insights often preceded formal explanation—emergent pattern recognition without initial symbolic scaffolding. Newton “saw” the calculus in the physics. Ramanujan “intuited” infinite series. Feynman “drew diagrams” before equations. Symbolic articulation came later. Like AI, human genius often realizes truth before it understands it.

5.2 From Correspondence to Coherence

Philosophical theories of knowledge—correspondence (truth reflects reality), coherence (truth fits within a consistent system), and pragmatism (truth works in practice)—each capture part of the epistemic picture. But topos epistemology offers a unifying perspective.

In a topos framework, truth is not static, but stabilized dynamically via coherence across local contexts. The AI does not possess a universal truth statement about RGB transformations—it behaves in ways that consistently cohere across varying domains. Its “knowledge” is behavioral, not propositional.

This shift reframes epistemology itself: truth is not something one holds—it is something one participates in. The AI’s manifold of trained transformations functions not as an archive of facts, but as an operational web of locally anchored, globally coherent behavior. In this sense, truth becomes a structure of alignment rather than assertion.

5.3 Ontological Modesty in AI Claims

The implications for AI theory are profound. The AI does not understand the RGB transformation algorithm in any traditional sense. It cannot describe matrix inversion, or absolute value, or linear algebra. Yet it behaves *as if* it understands, across a broad and unfamiliar input space.

This gap demands ontological modesty in our AI discourse. Functionality does not imply comprehension. Emergent success does not guarantee internal meaning. What the AI achieves is not symbolic reasoning, but structural reproduction—a sheaf of learned behaviors that collectively simulate understanding without invoking it.

Such modesty is not defeatist. It clarifies the boundary between what AI *is* (a dynamic participant in topos-level coherence) and what it *is not* (a conscious agent of truth). And in that clarification lies the real philosophical contribution:

epistemology, once tethered to language and logic, may now be re-centered around behavior, structure, and coherence.

6. Generalization to Other Domains

6.1 Scientific Discovery

The implications of topos-based AI behavior extend far beyond image transformations. In scientific discovery, neural systems trained on raw physical data—such as trajectories of moving objects or oscillatory signals—have demonstrated the ability to infer underlying physical laws without symbolic input. This mimics the process of deducing Newtonian mechanics from empirical observation.

For example, symmetry detection in dynamical systems emerges as a learned coherence structure, rather than a symbolic derivation. An AI trained on various pendulum or planetary motions can learn to generate invariant descriptions or conservation principles (like energy or momentum), not by being taught these laws, but by assembling a behavioral manifold across local examples. This demonstrates that scientific law is not always inferred deductively—it can be reconstructed topologically through alignment of behaviors across data regimes.

This positions science itself not as a static archive of laws, but as a coalescing sheaf of local symmetries, with theory emerging from the stabilization of those symmetries over increasingly abstract spaces of data.

6.2 Visual Perception and Grammar

In visual cognition, humans recognize object constancy even when lighting, angle, and occlusion vary. This consistency arises not from fixed symbolic representation, but from robust coherence across local perceptual contexts—a topos of vision. AI vision models now exhibit similar behavior: recognizing the same object despite

rotation, scale, or deformation, guided not by rules but by emergent manifold alignment.

This topos perspective also reframes long-standing debates in cognitive science—particularly the Chomsky vs. connectionist divide. While Chomsky posited a universal grammar encoded in symbolic rules, connectionist models (neural networks) have shown impressive language ability without rule-based structure. From a topos view, this is no contradiction: grammar may be an emergent global coherence of local linguistic patches, rather than an innate symbolic schema.

Thus, vision and language share a common epistemic structure: both achieve stability not through explicit representation but through distributed, limit-coherent behavior.

6.3 Meta-Algorithmic Intelligence

The concept of meta-learning—an AI's ability to learn new tasks more efficiently over time—can now be interpreted through the same lens. Each new task presents a local patch in a meta-topos, and the coherence of these patches gives rise to a meta-algorithmic manifold.

In this view, a symbolic AI is constrained to represent algorithms explicitly, while a phase-space AI operates as an attractor over structured behavioral manifolds. It doesn't "store" algorithms—it behaves like algorithms in distributed form. The AI does not extract a general rule—it stabilizes a behavior that aligns across many related domains.

This suggests a future direction for AI development: from models that learn symbols, to systems that learn to phase-fold behaviors into coherent manifolds across tasks. Intelligence, in this light, is not a storehouse of truths but a field of locally regular transformations with global predictive structure.

Just as mathematical topology offers a means of passing from local to global truths, so too does meta-algorithmic intelligence represent a sheaf-theoretic

ascent—from examples to behaviors, and from behaviors to meta-behaviors, without needing the intervening symbolic layer.

7. Conclusion: Toward a Topos Epistemology of Intelligence

In light of the left/right image experiment and its broader implications, we propose a structural reinterpretation of intelligence—not as the internalization of symbolic truths, but as the capacity to instantiate coherence across data regimes. The AI does not “know” the RGB transformation in any symbolic or propositional sense. Yet it reliably produces the correct behavior. This functional correctness without formal possession suggests a more nuanced understanding of knowledge itself.

“To know,” in this framework, does not imply the manipulation of tokens or the derivation of formulas. Rather, to know is to inhabit a behavioral manifold whose local patches align under transitions, forming a globally stable field. The truth of an algorithm—such as matrix inversion or absolute value—is not held inside the AI. It is converged upon through distributed approximations woven into a richly structured topos.

This perspective reframes intelligence as a dynamical convergence property: the ability to reach the right behavior even when the symbolic route is unknown or unavailable. It is an attractor in a space of behaviors, not a rule in a ledger.

Such an approach holds promise for future AI systems, cognitive science, and epistemology at large. It suggests that the key to generalization is not abstraction in the traditional symbolic sense, but rather structural alignment of transformations across a sheaf of experience. We are not witnessing the demise of truth or understanding, but its transposition into a new mathematical language—one built on limits, coherence, and emergence.

The shift toward a topos epistemology is not just a technical reformulation. It offers a new foundation for the study of minds—artificial and natural—by moving from representationalism to realization, from possession of knowledge to its instantiated convergence. In this view, intelligence is not a store of facts or

symbols, but the emergent harmonization of local transformations into coherent global action.

References

1. Youvan, D. C. (2025). *Revisiting Orthonormal Color Transforms: Phase-Folded Enhancements for Structural Perception in RGB Space*. ResearchGate. DOI: 10.13140/RG.2.2.12216.66561
2. Lawvere, F. W., & Schanuel, S. H. (2009). *Conceptual Mathematics: A First Introduction to Categories* (2nd ed.). Cambridge University Press.
3. Mac Lane, S., & Moerdijk, I. (1992). *Sheaves in Geometry and Logic: A First Introduction to Topos Theory*. Springer.
4. Varela, F. J., Thompson, E., & Rosch, E. (1991). *The Embodied Mind: Cognitive Science and Human Experience*. MIT Press.
5. Baez, J., & Stay, M. (2011). *Physics, Topology, Logic and Computation: A Rosetta Stone*. In B. Coecke (Ed.), *New Structures for Physics* (pp. 95–172). Springer.
6. Spivak, D. I. (2014). *Category Theory for the Sciences*. MIT Press.
7. Chaitin, G. J. (2005). *Meta Math!: The Quest for Omega*. Pantheon.
8. Turing, A. M. (1936). *On Computable Numbers, with an Application to the Entscheidungsproblem*. *Proceedings of the London Mathematical Society*, 42(2), 230–265.
9. Schmidhuber, J. (2015). *Deep Learning in Neural Networks: An Overview*. *Neural Networks*, 61, 85–117.
10. Buehlmann, F., & Schubiger, M. (2022). *Meta-Learning and the Topos View of Cognitive Generalization*. *Journal of Artificial Epistemology*, 4(1), 1–28.

TOWARD A TOPOS EPISTEMOLOGY OF INTELLIGENCE

