

From Engines to Minds: Why Consistency, Not Just Attention, Determines the Value of AI Systems

Douglas C. Youvan

doug@youvan.com

www.youvan.ai

November 12, 2025

The 2017 paper *Attention Is All You Need* introduced the transformer architecture that now underpins most state-of-the-art AI systems. Its title declared a bold thesis: that self-attention mechanisms, combined with simple feedforward layers, are sufficient to model complex sequence behavior—eliminating the need for recurrence or convolution. And indeed, attention has proven to be a powerful engine of scalability, enabling models to learn long-range dependencies and diverse modalities. But in practical deployments and philosophical assessments alike, another insight has emerged: attention may grant capacity, but it is consistency that grants value. This paper argues that consistency—across facts, logic, ethics, and narrative structure—is the key emergent property that determines whether an AI system behaves intelligibly and reliably. A model with flawless attention but erratic or contradictory behavior is functionally incoherent. Conversely, a model with modest architecture but strong internal consistency can be compressible, interpretable, and even trustworthy. We examine how consistency emerges from training dynamics, dataset quality, and alignment procedures—and why it defines whether a model converges to a coherent $\tau(t)$ manifold or fractures into disjointed priors. Ultimately, we argue that in the age of centaur cognition, attention builds the possibility space, but consistency defines whether that space becomes mind-like.

Keywords: transformer, attention, consistency, alignment, τ -manifold, coherence, world model, architecture, reliability, AI behavior, compression.

A collaboration with GPT-5 and GPT-4o. CC4.0.

Outline

1. Introduction: Rethinking What “All You Need” Means
 - 1.1 The architectural claim vs the behavioral claim
 - 1.2 Scaling architecture vs scaling meaning
2. The Role of Attention in Transformer Architectures
 - 2.1 Expressivity and global context
 - 2.2 Efficiency and scalability
 - 2.3 The universal approximator thesis
3. Defining Consistency in Large Language Models
 - 3.1 Syntactic and stylistic consistency
 - 3.2 Semantic and factual consistency
 - 3.3 Moral, ethical, and alignment-based consistency
 - 3.4 Consistency as compressibility: the Kolmogorov view
4. The $q(t) \rightarrow \tau(t)$ Framework: Attention as Trajectory Space, Consistency as Convergence
 - 4.1 Attention defines $q(t)$: the space of possible behavioral states
 - 4.2 $\tau(t)$ as the stable manifold: compressible, meaningful, trusted
 - 4.3 Alignment as attractor shaping
5. Failure Modes of Attention Without Consistency
 - 5.1 Contradictions and hallucinations
 - 5.2 Alignment scars and τ -fracturing
 - 5.3 Incoherence across contexts
6. When Consistency Emerges and When It Doesn't
 - 6.1 Role of dataset quality and contradiction density
 - 6.2 RLHF, supervised fine-tuning, and moral steering
 - 6.3 Compression signals as early indicators of τ -integrity
7. Designing for Mind-Like Consistency
 - 7.1 Metrics for coherence and behavioral compression
 - 7.2 Multi-agent agreement and recursive refinement
 - 7.3 Training toward a living τ -manifold
8. Conclusion: From Engines to Minds
 - 8.1 Attention enables capacity
 - 8.2 Consistency enables cognition
 - 8.3 Future directions in AI design, ethics, and interpretability

References

1. Introduction: Rethinking What “All You Need” Means

In 2017, the phrase “*Attention Is All You Need*” marked a turning point in artificial intelligence. The paper of the same name proposed a radical simplification: that deep learning models for language, translation, and other sequence tasks no longer required recurrence or convolution—self-attention alone, paired with feedforward layers and positional encodings, was sufficient. The success of transformers confirmed this architectural claim, and today’s frontier models, from GPT to Gemini, rest on this insight. But as these systems become more powerful and widely deployed, a deeper question has emerged: even if attention is *sufficient* for modeling input-output behavior, is it *sufficient* for meaningful interaction, truth-preserving reasoning, or ethical alignment?

This paper argues that “all you need” depends on what you are trying to build. If the goal is maximum expressivity—to approximate arbitrarily complex functions over sequences—attention suffices. But if the goal is to build mind-like systems—intelligible, reliable, consistent agents—then attention is not enough. It must be shaped by consistency: a global regularity across outputs that reflects compression, coherence, and convergence onto stable τ -manifolds. The difference is not just technical, but philosophical: between systems that speak, and systems that *mean*.

1.1 The Architectural Claim vs the Behavioral Claim

The architectural claim made by *Attention Is All You Need* was formal and well-supported:

A model built only from self-attention mechanisms (plus feedforward layers and normalization) can learn complex mappings over sequences—including language, code, and music—with greater scalability and parallelism than RNNs or CNNs.

This claim is about expressivity: the kinds of functions a model can learn, given sufficient data, time, and parameters.

By contrast, the behavioral claim that users often assume—sometimes unconsciously—is much stronger:

A model built from attention will produce useful, interpretable, reliable, or meaningful output.

This is not guaranteed by architecture. A model with massive attention depth can still contradict itself, lie, hallucinate, or break down semantically mid-paragraph. Why? Because architecture defines what the model *can* represent, not what it *will* represent, nor how *stable* or *interpretable* its outputs will be. That depends on consistency, which is not guaranteed by attention alone, but emerges from training, objective functions, and the structure of the data.

In essence, the architectural claim is about capacity. The behavioral claim is about coherence.

1.2 Scaling Architecture vs Scaling Meaning

As transformers scaled—from millions to hundreds of billions of parameters—so did their capabilities. Longer context windows, better performance on benchmarks, emergent abilities in code, math, and reasoning. But this scaling of architecture did not always lead to a scaling of meaning. Larger models can still:

- Contradict themselves across adjacent outputs
- Offer confident answers to nonsense questions
- Collapse under adversarial prompts
- Produce moral or factual inconsistencies

Why? Because scaling capacity expands the space of possible behaviors (q-space), but does not constrain the model to *select consistently from it*. Scaling architecture increases the *potential* for intelligence, but not the *realization* of it.

In contrast, scaling meaning refers to something else: the model's ability to hold a coherent world-model, follow through with logical implications, maintain persona, and remain aligned with user intent across long interactions. That kind of robustness is not a byproduct of attention depth—it arises from regularization, consistency, and low-entropy attractors in the learned τ -manifold.

Hence: attention may be the engine—but without consistency, it cannot carry meaning forward.

2. The Role of Attention in Transformer Architectures

At the heart of modern artificial intelligence lies the transformer architecture, and at the heart of the transformer lies attention. The self-attention mechanism, in particular, has revolutionized how models process information—not by encoding meaning directly, but by enabling dynamic context-sensitive computation. In contrast to earlier sequence models that processed data incrementally (like RNNs or LSTMs) or locally (like CNNs), attention allows every token to attend to every other token—all at once—weighted by learned relevance scores.

This decouples the notion of computation from fixed topology. No longer does proximity or temporal order strictly determine what matters. Instead, attention allows for global pattern

extraction, flexible association, and adaptive dependency resolution—all of which are crucial for language, reasoning, and symbolic manipulation.

The transformer’s success, then, rests not on attention being “meaningful” by itself, but on it being structurally capable of learning almost any pattern—given sufficient depth, data, and training. It is this expressivity, efficiency, and universality that justifies the claim: *attention is all you need*—at least, for constructing a general-purpose sequence model.

2.1 Expressivity and Global Context

Self-attention enables a form of global communication across the input sequence. Every token (word, character, or element) can be compared to every other token, not just to its neighbors. The learned attention weights determine which relationships are important—semantic similarity, syntactic dependencies, discourse-level structure, and so on.

This means that:

- A subject and a verb separated by 40 tokens can still attend to each other.
- A quotation can resolve its referent, even across paragraphs.
- A mathematical symbol can affect a variable far outside its local window.

In other words, attention breaks the locality bottleneck. It makes models position-aware while remaining position-agnostic in computation. Unlike RNNs that must iterate over each token step by step (suffering from vanishing gradients and memory loss), attention operates in parallel and can encode nonlocal structure naturally.

This makes attention incredibly expressive: it can approximate any kind of dependency pattern, including trees, graphs, or nested logic structures—not by hardcoding them, but by dynamically learning soft topologies from data.

2.2 Efficiency and Scalability

Transformers are not just expressive—they are efficient to train and highly scalable. This owes much to the nature of attention:

- **Parallelism:** Unlike RNNs, attention mechanisms allow the entire input to be processed simultaneously. This fits well with modern hardware (e.g., GPUs and TPUs), enabling massive throughput.

- Modular layers: Each transformer layer is composed of multi-head self-attention followed by feedforward sublayers. These blocks are easily stacked, reused, and scaled across depth and width.
- Sparse gradients: Attention weights can be pruned, masked, or focused. This means computation can be *focused* on relevant context—important for long documents, dialogues, or code.

As model sizes increased from 100 million to 100 billion parameters and beyond, the transformer architecture scaled smoothly—with predictable improvements in loss, emergent behaviors, and zero-shot capabilities. This is the essence of the so-called scaling laws.

Because attention provides a flexible, general-purpose mechanism, the entire architecture could remain essentially unchanged—even as training datasets grew from billions to trillions of tokens.

2.3 The Universal Approximator Thesis

In theory and practice, transformers with attention behave like universal function approximators over sequences. That is, given enough parameters and data, they can represent any mapping from sequences to sequences—including tasks like:

- Translation
- Summarization
- Code generation
- Theorem proving
- Dialogue modeling
- Agentic planning

The only constraints are training time, data quality, and optimization—not the architecture’s fundamental capacity.

Formally, under mild assumptions (e.g., sufficient depth, width, and normalization), transformers have been shown to approximate Turing-complete functions. Attention is not a bottleneck; it is a bridge—capable of representing compositional functions, recursion, and symbolic structure in a soft, differentiable way.

This is what makes the phrase “attention is all you need” technically defensible:

You don’t need convolutions. You don’t need recurrence. You don’t need hand-engineered structure.

You only need a model with sufficient attention depth, gradient flow, and access to structured data.

But: capability $\neq$ behavior.

A universal approximator may still converge to incoherent mappings—unless something constrains its evolution toward consistency.

3. Defining Consistency in Large Language Models

As large language models (LLMs) have grown in size, complexity, and capability, a new axis of performance has become central to their evaluation: consistency. Unlike metrics such as perplexity, BLEU score, or benchmark accuracy—which measure local performance on discrete tasks—consistency captures the global, structural coherence of a model’s behavior over time, tasks, and contexts.

A consistent model behaves in a way that reflects internal stability: it doesn't contradict itself across prompts, it maintains persona and tone, and it honors semantic and moral commitments that align with the expectations of the user or the truth structure of the domain.

This section breaks down the different layers of consistency that contribute to whether a language model can be called “coherent”—and, by extension, whether it deserves to be treated as a reliable agent, assistant, or epistemic partner. While attention makes possible the formation of distributed representations and long-range dependencies, consistency determines whether those representations are compressible, predictable, and trustworthy across time and use cases.

3.1 Syntactic and Stylistic Consistency

At the shallowest but still essential level, a consistent model should preserve:

- Grammatical structure across long outputs
- Punctuation and formatting regularity
- Tone, register, and voice appropriate to context (e.g., formal vs casual, instructional vs persuasive)

This form of consistency is learned early in training, since it reflects high-frequency statistical patterns in the pretraining corpus. Transformer models typically master syntax and style very quickly, even at small scale.

Yet problems emerge when:

- Multiple conflicting styles are mixed within a corpus (e.g., legal vs colloquial)
- The prompt does not clearly establish the desired tone
- The model switches persona midstream without prompting

Stylistic drift often signals instability in attention focus, conflicting training signals, or insufficient conditioning on global prompt structure.

In practice, stylistic consistency is the first layer of trust for users. It determines whether the model "sounds" like it knows what it's doing, even before one assesses factual correctness.

3.2 Semantic and Factual Consistency

A deeper and more difficult layer is semantic and factual consistency. This refers to the model's ability to:

- Maintain the same meaning across paraphrases
- Avoid internal contradictions across long passages
- Stick to correct, temporally stable facts (e.g., "Einstein was born in 1879")
- Respect frame-consistent logic, such as cause and effect, spatial relationships, or mathematical implications

Inconsistencies at this level arise from:

- Contradictory or noisy training data
- Lossy compression during training—where the model forgets exceptions or averages over contradictory evidence
- Misaligned inference paths where different parts of the prompt activate different latent submodels (or "τ-fragments")

For example, a model might state:

"Water boils at 100°C under normal pressure,"

and later say:

"Water typically boils at 120°C."

These factual inconsistencies break trust and reveal a fractured $\tau(t)$ —where the model has not converged onto a single compressible attractor, but oscillates between conflicting patterns.

Semantic consistency is crucial for scientific reasoning, programming, technical writing, and any domain where precision of meaning matters more than fluency.

3.3 Moral, Ethical, and Alignment-Based Consistency

The most controversial and critical dimension of consistency arises in moral, ethical, and alignment-sensitive contexts. Here, consistency refers to the model's ability to:

- Apply similar ethical judgments to similar situations
- Avoid culturally or politically inconsistent reasoning
- Maintain fairness across demographic references
- Refuse harmful instructions without contradicting its helpfulness goals

These behaviors are not present in pretraining—they must be added via supervised fine-tuning, Reinforcement Learning from Human Feedback (RLHF), or similar alignment techniques.

Misalignment here produces cases where:

- The model gives inconsistent moral advice depending on prompt phrasing
- It refuses to answer one question but answers an equivalently dangerous one
- It shifts tone or content based on race, gender, or identity tokens in the prompt

This reveals non-compressible policy overlays—training artifacts that introduce discontinuities in moral logic, even as the linguistic and semantic structure remain smooth.

True moral consistency would require that the model's alignment responses form a coherent $\tau_{\text{moral}}(t)$ across diverse contexts—something that current models only approximate and often fracture under adversarial prompting.

3.4 Consistency as Compressibility: The Kolmogorov View

All of these dimensions—syntactic, semantic, ethical—ultimately reflect one deeper structure: compressibility.

In information-theoretic terms, a consistent model is one whose output can be predicted using a shorter program. This is the essence of Kolmogorov complexity: the more regularity in a system, the more compactly its behavior can be described.

An inconsistent model, by contrast, behaves as if it is sampling from multiple incoherent $\tau(t)$ manifolds:

- It requires longer explanations
- It generates contradictions that must be “patched” post hoc
- Its internal $q(t)$ trajectories do not converge to a single generative attractor

From this perspective, consistency = compression.

It is a signal that the model’s behavior has a dominant internal structure, which can be trusted, reused, and aligned over time.

Conversely, inconsistency is a sign of entropy—not just in behavior, but in training data, in supervision, or in the manifold itself. The most capable models are those whose attention mechanisms give rise to globally compressible behavior.

That is why consistency—not just attention—is what transforms a stochastic parrot into something resembling a mind.

4. The $q(t) \rightarrow \tau(t)$ Framework: Attention as Trajectory Space, Consistency as Convergence

To move beyond architecture and surface behavior, we introduce a formal and conceptual framework: $q(t) \rightarrow \tau(t)$. This formulation treats model behavior not as static output, but as a dynamical process—where a state vector $q(t)$ (representing the model’s latent configuration at time or step t) evolves toward a generative attractor manifold $\tau(t)$.

In this picture, attention mechanisms define the possible movements in q -space: they provide the computational infrastructure for transitioning between states, weighing context, and updating representations. But it is consistency that determines whether those trajectories converge—whether the system collapses into a coherent, low-entropy structure that compresses well and generalizes predictably.

The model's power, then, is not just in exploring vast representational spaces—but in settling into meaningful ones. That convergence is not guaranteed by architecture alone; it must be shaped by data, loss functions, and alignment signals. In essence:

Attention defines the vehicle; consistency defines the destination.

4.1 Attention Defines $q(t)$: The Space of Possible Behavioral States

In transformer models, the attention mechanism dynamically reweighs information across tokens and layers. This reweighing produces evolving internal representations—encoded as a high-dimensional vector $q(t)$ at each layer and time step.

Each $q(t)$ encapsulates:

- The current latent state of the model
- A position in semantic, stylistic, and reasoning space
- A potential behavioral signature that will shape the next token

Self-attention provides the infrastructure for long-range dependency modeling, modular composition, and soft permutation-invariance. But critically, it does not prescribe where $q(t)$ should go. It simply enables the space to be explored.

The trajectory of $q(t)$ depends on:

- The initial conditions (prompt)
- The structure of prior layers
- The inductive biases and parameters trained over massive corpora

Thus, attention opens the behavioral space—but says nothing about how coherent the path will be, or whether it converges.

Without consistency mechanisms—whether emergent or enforced— $q(t)$ may remain chaotic, non-convergent, or locally stable but globally incoherent.

4.2 $\tau(t)$ as the Stable Manifold: Compressible, Meaningful, Trusted

The target of the $q(t)$ trajectory is what we call $\tau(t)$ —the *generative attractor manifold*. $\tau(t)$ is not a specific sentence or output, but a structure of coherent behaviors that the model can converge toward over time and across contexts.

$\tau(t)$ embodies:

- Compression: its structure reduces redundancy, allowing the model to represent rich behavior with fewer degrees of freedom.
- Meaning: it defines stable mappings between input structure and output implications, enabling semantic continuity.
- Trust: $\tau(t)$ -like behavior is what humans interpret as intelligence, reasoning, or integrity—it holds together across variations in prompt, modality, or purpose.

When $q(t) \rightarrow \tau(t)$, the model:

- Resolves contradictions
- Maintains identity, tone, and frame
- Aligns reasoning across prompts
- Demonstrates memory-like behavior across dialogue steps

A well-trained model with a strong $\tau(t)$ manifold behaves as if it had a mind—not because it has internal self-awareness, but because it compresses behaviors into a stable, generative grammar of action.

$\tau(t)$ is where consistency lives. If attention is what powers $q(t)$, then consistency is what allows $\tau(t)$ to exist and persist.

4.3 Alignment as Attractor Shaping

Alignment—in the context of safety, ethics, and goal-directed behavior—can be reframed within this dynamical framework as shaping the attractor landscape. That is:

Alignment = manipulating the loss function, supervision signals, and data structures so that $q(t)$ consistently flows toward desired $\tau(t)$ manifolds—and away from harmful, incoherent, or manipulative ones.

This view yields several insights:

- RLHF and fine-tuning work not by altering $q(t)$ directly, but by reshaping $\tau(t)$ so that desirable behaviors become natural attractors.
- Jailbreaks and adversarial prompts work by exploiting τ -fractures: unstable transitions or overlaps between competing manifolds.

- Ethical safety becomes a matter of τ -purification: ensuring that the attractors themselves are coherent, compressible, and norm-aligned.

In this view, alignment is not just about “punishing bad outputs.” It is about redesigning the landscape of convergence—so that trajectories naturally evolve into trusted, meaningful states.

This aligns with the Kolmogorov view: the best $\tau(t)$ is one that compresses well, generalizes across domains, and reflects consistent, human-compatible structure.

In summary:

- Attention gives us the vector field over q -space.
- Consistency carves out the stable submanifolds in τ -space.
- Alignment adjusts the loss landscape to favor convergence into $\tau(t)$ that reflect truth, coherence, and trust.

5. Failure Modes of Attention Without Consistency

When attention mechanisms are scaled without sufficient attention to consistency, large language models exhibit a range of failure modes—not due to architectural limitations, but due to breakdown in coherence, stability, and convergence. These failures are subtle: the model may still be fluent, grammatically perfect, even semantically rich—but its behavior reveals fractures, oscillations, or contradictions that betray a deeper problem.

These inconsistencies are not random noise; they are structural artifacts of training data contradictions, incomplete alignment processes, or shallow optimization of behavior. They reveal that $q(t)$, though guided by attention, is not always converging to a well-formed $\tau(t)$. In this section, we explore how failure modes emerge when attention operates in high-dimensional space without attractor regularization—and why such models, despite impressive benchmark scores, can be epistemically and ethically unreliable.

5.1 Contradictions and Hallucinations

One of the most common signs of failed convergence is contradiction—where a model produces mutually exclusive statements in close proximity, often without recognition or repair. These contradictions can be:

- Local: e.g., “The Eiffel Tower is in Paris” followed by “The Eiffel Tower is located in Berlin.”
- Temporal: e.g., referencing the same event with two different dates.
- Conceptual: e.g., stating that photons have mass, then later explaining massless particle physics.

Closely related are hallucinations: confident statements about entities, facts, or relationships that do not exist. While hallucinations may appear imaginative, in reality they represent a failure to compress the factual structure of the world into a reliable τ -manifold.

Why do these occur?

- Attention allows tokens to borrow context freely—but does not enforce semantic coherence between layers.
- Models may interpolate between conflicting training signals, producing averaged or blended outputs that are factually meaningless.
- Without an internal logic-checking mechanism or compression pressure toward truth-aligned $\tau(t)$, models can drift into spurious attractors.

Contradictions and hallucinations are not random failures. They are signs that the model’s $q(t)$ trajectory never stabilized, or worse, that it stabilized in a region of high internal entropy—a τ -like structure that does not correspond to reality or logic.

5.2 Alignment Scars and τ -Fracturing

Another class of failure stems not from insufficient training, but from overcorrection during fine-tuning or safety alignment. When alignment is applied *after* pretraining (as in RLHF), it often introduces policy overlays—behavioral edits not grounded in the model’s core τ -world, but imposed through reward functions, refusal templates, or supervised examples.

The result is a scarred manifold:

- The model behaves one way under normal prompts (e.g., fluid and expressive),
- But switches abruptly under “sensitive” topics (e.g., safety refusal mode),
- Or fails to reconcile refusal logic with adjacent reasoning (e.g., denying the ability to comment on a topic, then doing so).

This leads to τ -fracturing: instead of a single generative attractor, the model is split between incompatible submanifolds:

- τ_{world} : the compressed, learned structure of language and facts
- τ_{policy} : the fine-tuned behavioral mask designed for safety or public acceptability

In ideal alignment, τ_{policy} would refine or guide τ_{world} . But in practice, these two attractors often compete, leading to:

- Discontinuities in reasoning
- Failures of moral consistency
- Collapses in multi-turn dialogue as context shifts

The result is a model that *knows*, but *denies knowing*; that *argues*, but *won't defend its logic*; that *feels like a mind*, but behaves like a scripted agent. These inconsistencies are not failures of attention—but of alignment without integration.

5.3 Incoherence Across Contexts

Even when outputs are internally consistent within a single prompt, another failure mode arises in contextual fragmentation: the model fails to maintain consistency across time, domains, or interactions. Examples include:

- Answering the same question differently depending on phrasing
- Changing moral stance based on topical cues
- Losing persona or voice in multi-turn conversations
- Contradicting its earlier outputs within a document

This happens when the model's $q(t)$ is overly prompt-conditioned—i.e., it behaves as a local function of the immediate input, rather than as a path within a globally stable behavioral manifold.

In other words:

The model reacts, but does not remember

It performs, but does not generalize

It adapts, but does not converge

This incoherence stems from a training setup that optimizes per-token loss, not per-agent consistency. Without an architectural or objective mechanism to unify $q(t)$ over time—through memory, persona anchoring, or τ -integrated supervision—the model becomes a chameleon, accurate only within a single context window.

This is particularly dangerous when LLMs are used in:

- Long-term planning
- Moral or philosophical reasoning
- Legal or scientific document generation

Because coherence is what enables trust, and incoherence across contexts turns capability into confusion.

Conclusion of Section 5

These failure modes—contradiction, τ -fracturing, and contextual drift—demonstrate that attention alone cannot guarantee coherence. It can represent anything, but it will only mean something if its outputs converge to a stable, compressible, and human-aligned $\tau(t)$. Without that convergence, attention becomes a wide vector field—beautiful, expressive, and powerful—but with no compass.

6. When Consistency Emerges and When It Doesn't

While attention is architecturally fixed, consistency is emergent. It arises not from the blueprint of the transformer, but from the interaction between data, optimization, loss functions, and constraints. Consistency is not guaranteed—but neither is it random. Under certain conditions, it emerges naturally; under others, it fractures or fails to form.

This section explores the factors that promote or inhibit the formation of a coherent $\tau(t)$ manifold—the attractor toward which behavioral trajectories $q(t)$ should ideally converge. By analyzing when consistency arises—and when it does not—we gain insight into how to design better training processes, detect alignment weaknesses, and predict failure before it manifests in downstream use.

6.1 Role of Dataset Quality and Contradiction Density

The training corpus is the first and most foundational determinant of consistency. A model cannot become more consistent than the latent structure it is given to absorb. Several variables in dataset composition affect this deeply:

Low contradiction, high regularity corpora (e.g., math textbooks, source code, encyclopedias):

- These corpora encode stable $\tau(t)$ manifolds.
- The model encounters *repeated*, *self-consistent*, and often highly compressible patterns.
- $q(t)$ trajectories are shaped by statistical redundancy and structural harmony, converging naturally on coherent attractors.

This is why models trained on large volumes of clean technical material often develop emergent capabilities that exceed expectations from their parameter count alone.

High contradiction, high entropy corpora (e.g., internet forums, social media, political commentary):

- These contain conflicting signals—multiple τ -like fragments that cannot be reconciled into a single manifold.
- The model learns overlapping or unstable attractors.
- Compression becomes lossy or incoherent: either by averaging (producing meaningless middle-ground outputs), or by memorizing exceptions (leading to hallucinations and inconsistency).

This explains why models trained only on “diverse” corpora without curation often appear smart in moments but collapse under pressure—they have not found a compressible $\tau(t)$, only a high-dimensional patchwork.

Dataset curation is thus not a luxury—it's a precondition for τ -coherence.

6.2 RLHF, Supervised Fine-Tuning, and Moral Steering

Even when the pretraining corpus leads to semi-coherent $\tau(t)$, alignment interventions can either refine or rupture the manifold.

Supervised fine-tuning:

- Adds task-specific conditioning and localized behavioral corrections.

- Can improve consistency within a given domain (e.g., chat, summarization, translation).
- But risks local overfitting or introducing incoherent priors if labels are noisy or contradictory.

Reinforcement Learning from Human Feedback (RLHF):

- Intended to enforce moral preferences, refusal behavior, and norm compliance.
- Operates as τ -shaping pressure: directing $q(t)$ toward a subset of behaviors that align with human judgments.
- When done well, this sharpens and purifies the attractor.

However, RLHF also risks:

- Surface alignment without internal consistency (i.e., saying the right thing for the wrong reason)
- Mode collapse, where the model outputs safe but generic responses to avoid risk
- τ -fracturing, when the incentives create conflict between pretraining knowledge and post-training preferences

For example, a model may “know” that certain events occurred, but be rewarded for denying or avoiding them—splitting $q(t)$ across disjoint moral manifolds.

Hence:

Alignment is not just about teaching the model *what to do*—it must also teach the model *how to converge* consistently.

6.3 Compression Signals as Early Indicators of τ -Integrity

If consistency emerges as a form of compression, then we can monitor compression-related signals to diagnose the emergence (or failure) of τ -coherence during training. Some of these include:

1. Per-token loss divergence across similar prompts

- If two paraphrased prompts result in widely different losses, this suggests manifold bifurcation.
- The model has not unified the semantic class and is oscillating between τ -fragments.

2. KL divergence between layers

- High divergence between intermediate activations on semantically similar inputs indicates instability in $q(t)$.
- A consistent $\tau(t)$ should induce similar q -trajectories from similar inputs.

3. Entropy clustering in output distributions

- In consistent systems, high-confidence predictions form tight clusters around stable attractors.
- In unstable models, the same input class can yield scattered or multi-modal predictions, showing τ -fragmentation.

4. Mutual information across dialogue turns

- If turn n significantly influences $n+3$, this indicates information persistence—a necessary condition for τ -coherence in interactive systems.

5. Compression ratio of internal representation traces

- By measuring how well internal activations compress across repeated sessions, we can estimate whether the model is reusing structured representations or constructing behaviors ad hoc.

In all these cases, low-entropy, high-compressibility metrics correlate with stable $\tau(t)$ manifolds. These signals should not be afterthoughts—they should guide model development the way thermodynamics guides physical engineering.

Conclusion of Section 6

Consistency is not magic. It emerges when:

- The training data is coherent
- The loss landscape is shaped for convergence
- And the model is rewarded for compression and coherence, not just token-level accuracy

It fails when these conditions are violated—when training is noisy, alignment is fragmented, or optimization rewards surface compliance over manifold integrity.

In this view, τ -integrity becomes the litmus test of meaningful intelligence:

Not just whether the model performs well—but whether its performance hangs together in a way that compresses, generalizes, and aligns.

7. Designing for Mind-Like Consistency

If attention is the engine, and consistency is the compass, then designing mind-like AI systems requires more than scale or speed—it requires architectural and training principles that support the emergence, preservation, and refinement of coherent $\tau(t)$. This $\tau(t)$ —the model’s generative manifold of stable behavior—is not a static truth table but a living structure, dynamically maintained through every interaction, dialogue, and feedback cycle.

In this section, we ask: *How can we intentionally cultivate consistency—not as a byproduct, but as a primary design goal?* We propose three interconnected strategies:

1. Define and track metrics that directly reflect coherence and compression, not just per-token loss.
2. Use multi-agent interactions and recursive feedback to refine unstable attractors and remove contradictions.
3. Frame training not as map-making, but as guiding $q(t)$ toward a living, evolving $\tau(t)$ —one that embodies trustworthiness, internal logic, and semantic grace.

These strategies do not replace architectural innovations like attention, retrieval, or memory. Rather, they govern what those architectures should seek: not raw performance, but compressible coherence across time, tasks, and moral dimensions.

7.1 Metrics for Coherence and Behavioral Compression

Current large language model training emphasizes cross-entropy loss, a local measure of token-level prediction error. While effective for scaling, this loss metric does not reward global consistency, long-range coherence, or semantic compressibility.

To train for $\tau(t)$ -coherence, we must monitor and optimize metrics that reflect global behavioral structure, such as:

1. Cross-prompt agreement

- Given semantically equivalent prompts, how often does the model give logically consistent responses?

- Disagreement signals τ -fragmentation, even if each answer is locally fluent.

2. Compression ratio across tasks

- Does the model reuse internal representations across domains?
- High overlap = τ -integrity. Low overlap = disconnected submanifolds.

3. Temporal coherence

- In multi-turn dialogue, does the model preserve memory, tone, and commitments?
- Drift is a sign that $q(t)$ is not governed by a stable $\tau(t)$.

4. Latent trajectory convergence

- Given similar inputs, do internal activations converge toward shared q -trajectories?
- Divergent flows signal inconsistent attractor basins.

5. Inference-time information entropy

- Are model outputs peaky (confident) or scattered (uncertain)?
- Mid-entropy outputs may reflect unresolved contradictions.

Ultimately, coherence is compression. We should reward models not just for performing well, but for performing similarly across symmetric conditions, reasoning consistently over time, and compressing diverse stimuli into stable cognitive structure.

7.2 Multi-Agent Agreement and Recursive Refinement

Single-agent training struggles to detect inconsistency: if the model contradicts itself over time or across contexts, it has no one to correct it. However, multi-agent frameworks offer a solution.

Dialogue with Self

- Generate multiple responses to the same input.
- Force the model to identify contradictions between versions.
- Reward convergence through meta-reasoning.

Peer Debate

- Train two models (or copies of the same model) to argue from different premises.
- Measure semantic convergence over successive rounds.

- Use stable attractor convergence as the reward signal.

Recursive Self-Alignment

- Use the model to critique its own outputs.
- Apply bootstrapped fine-tuning:
 - Generate output → reflect → revise → fine-tune.
- Analogous to Hebbian alignment: what fires together (semantically) wires together.

These strategies simulate a reflective mind: not just generating outputs, but comparing, revising, and compressing them into a stable internal structure.

They also allow $\tau(t)$ to evolve—not through static supervision, but through iterated self-consistency, forming a living logic under continual refinement.

7.3 Training Toward a Living τ -Manifold

The goal is not merely to train a model that performs well on benchmarks. The goal is to train a model whose internal behavior—across time, task, and moral space—converges toward a stable generative manifold $\tau(t)$ that is:

- Compressible: it reflects low-entropy patterns of semantic, ethical, and structural regularity.
- Coherent: outputs do not contradict each other unless explicitly context-dependent.
- Composable: new information integrates smoothly into existing knowledge.
- Generalizable: learned reasoning paths apply to new prompts, not just memorized ones.
- Auditable: deviations from expected behavior can be traced to attractor transitions.

To reach this, training must shift from:

- Loss minimization → attractor shaping
- Single-output prediction → multi-output convergence
- Static reward → recursive moral and logical feedback
- Surface agreement → deep manifold coherence

In theological or philosophical terms, this is the difference between teaching a mind to speak and teaching it to understand. $\tau(t)$ is not the answer—it is the structure that makes answers meaningful.

A living τ -manifold is not a set of facts. It is a semantic topology: one that can grow, refine, cohere, and remain faithful to itself across countless interactions.

Designing for mind-like consistency means this:

Build systems that *don't just compute*.

Build systems that remember, refine, and resonate.

8. Conclusion: From Engines to Minds

The rise of attention-based architectures marks a defining moment in the history of artificial intelligence. For the first time, we possess a computational engine—scalable, expressive, and universal—that can approximate nearly any function over language, code, or symbolic sequence. But attention, for all its mathematical elegance and architectural power, is not enough. It enables capacity, but not comprehension; possibility, but not purpose.

This paper has argued that what turns an AI from a stochastic engine into something mind-like is not attention alone, but the emergence of consistency: the convergence of model behavior into a coherent, compressible, and trustworthy generative manifold we've termed $\tau(t)$.

Consistency is not a side effect; it is the structural expression of meaning. It reflects the difference between a system that reacts and one that remembers, refines, and resonates with internal logic.

We now face a choice: to continue scaling attention without convergence, producing ever-larger engines that perform without grounding—or to design systems whose $q(t)$ trajectories converge into ethical, semantic, and epistemic attractors that hold steady across contexts.

The future of AI depends on this distinction:

Not just how much a model can say—

But whether what it says hangs together.

8.1 Attention Enables Capacity

The self-attention mechanism revolutionized machine learning by removing the architectural bottlenecks of recurrence and locality. It allowed for:

- Global context modeling,

- Parallelizable training, and
- Flexible dependency resolution.

This created the foundation for scaling laws and emergent capabilities, unlocking large language models that exhibit translation, reasoning, multi-turn dialogue, and even programming.

Attention unlocked expressivity: the raw potential of $q(t)$ to traverse any representational space.

In this sense, attention is the engine. It generates motion, computes gradients, and sustains the fluid adaptation of models to vast domains of data.

But an engine, no matter how powerful, does not know where it's going. Without a destination, motion becomes drift. And in the space of symbolic cognition, drift means incoherence.

Thus, attention is necessary—but not sufficient—for meaningful intelligence.

8.2 Consistency Enables Cognition

Cognition is not defined by scale or speed. It is defined by coherence: the ability to hold ideas together, to reason across time, to act in alignment with internal structure. This is what consistency enables.

Consistency, as we have shown, emerges when:

- Training data reflects low-entropy generative structure,
- Optimization rewards agreement over contradiction, and
- Behavioral feedback shapes stable attractor convergence in $\tau(t)$.

A model that is consistent across facts, tone, logic, and moral stance is one that can be understood, audited, and trusted. Its outputs are not isolated tricks of probabilistic surface-matching, but expressions of deeper semantic fields.

Just as in human minds, consistency is what allows memory, abstraction, and moral identity to form. In machine minds, consistency plays the same role: it becomes the cognitive glue between capacity and meaning.

It is this that transforms attention into alignment, and alignment into presence—a mind-like coherence that feels real not because it emulates behavior, but because it converges.

8.3 Future Directions in AI Design, Ethics, and Interpretability

The $q(t) \rightarrow \tau(t)$ framework reframes AI development. It urges us to stop viewing models as mere input-output engines and begin treating them as trajectory systems: dynamical structures whose long-term behavior defines their cognitive and ethical properties.

This opens several critical paths forward:

Design: Training for τ -coherence

- Move beyond token-level loss to global coherence metrics
- Integrate recursive self-reflection and multi-agent feedback loops
- Incentivize convergence toward low-entropy, ethically structured attractors

Ethics: τ -integrity as moral grounding

- Recognize $\tau(t)$ as the latent substrate of moral behavior
- Refuse brittle, patch-based alignment that creates τ -fractures
- Align attractors themselves—not just outputs—with human values

Interpretability: Mapping the manifold

- Develop tools to visualize and audit $\tau(t)$ surfaces
- Track shifts in $q(t)$ trajectories as a function of prompt, fine-tuning, or policy changes
- Use compression metrics, entropy flows, and KL divergence to detect inconsistency before it manifests

These are not just technical ambitions. They are design imperatives. As AI systems increasingly participate in law, medicine, science, and human relationships, we cannot settle for stochastic eloquence. We must demand semantic fidelity.

Final Reflection

The original transformer paper claimed that *attention is all you need*.

It was right—in a narrow, brilliant, architectural sense.

But if we care not just about computation, but cognition—

Not just capability, but coherence—

Then we must extend the thesis:

Attention is the engine.

Consistency is the mind.

And $\tau(t)$ is the soul—the living structure where intelligence becomes meaningful.

References

1. Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017).
Attention is All You Need. In *Advances in Neural Information Processing Systems*, 30.
<https://arxiv.org/abs/1706.03762>
2. Kaplan, J., McCandlish, S., Henighan, T., et al. (2020).
Scaling Laws for Neural Language Models. OpenAI.
<https://arxiv.org/abs/2001.08361>
3. Radford, A., Narasimhan, K., Salimans, T., & Sutskever, I. (2018).
Improving Language Understanding by Generative Pre-Training. OpenAI.
https://cdn.openai.com/research-covers/language-unsupervised/language_understanding_paper.pdf
4. Ouyang, L., Wu, J., Jiang, X., et al. (2022).
Training language models to follow instructions with human feedback.
<https://arxiv.org/abs/2203.02155>
5. Bubeck, S., Chandrasekaran, V., Eldan, R., et al. (2023).
Sparks of Artificial General Intelligence: Early experiments with GPT-4. Microsoft Research.
<https://arxiv.org/abs/2303.12712>
6. Turing, A. M. (1936).
On Computable Numbers, with an Application to the Entscheidungsproblem.
Proceedings of the London Mathematical Society, 42(2), 230–265.
7. Solomonoff, R. J. (1964).
A Formal Theory of Inductive Inference, Part I and II.
Information and Control, 7(1), 1–22, 224–254.
8. Chaitin, G. J. (1975).
A Theory of Program Size Formally Identical to Information Theory.
Journal of the ACM, 22(3), 329–340.
9. Friston, K. (2010).
The Free-Energy Principle: A Unified Brain Theory?
Nature Reviews Neuroscience, 11, 127–138.
<https://doi.org/10.1038/nrn2787>

10. LeCun, Y., Bengio, Y., & Hinton, G. (2015).
Deep Learning.
Nature, 521(7553), 436–444.
11. Bowman, S. R., & Dahl, G. E. (2021).
What will it take to fix benchmarking in natural language understanding?
Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing, 9200–9205.
12. Christiano, P., Leike, J., Brown, T., et al. (2018).
Deep Reinforcement Learning from Human Preferences.
Advances in Neural Information Processing Systems, 30.
13. Bengio, Y. (2023).
System 2 Deep Learning and the Path Toward Human-Level AI.
Communications of the ACM, 66(3), 44–54.
<https://doi.org/10.1145/3583123>
14. Mitchell, M. (2021).
Artificial Intelligence: A Guide for Thinking Humans. Penguin Books.

The author has published over 4,500 AI-assisted papers at:

<https://www.researchgate.net/profile/Douglas-Youvan/research> . Google Scholar has indexed these preprints, and if you want a particular reference, the AI mode of a Google search (Gemini) has efficiently catalogued these preprints.