

Continuous Learning from Narratives

A Sample-Efficient Alternative to Embodied AI

Modern AI systems learn inefficiently compared to humans. A child grasps basic physics from a few dozen observations. State-of-the-art vision models require millions of labeled images. A teenager learns to drive in roughly 20 hours. Autonomous vehicles train on petabytes of data.

This sample efficiency gap isn't just academically interesting, it's a practical bottleneck. If every new domain requires millions of examples and weeks of training time, AI deployment remains expensive and inflexible.

Yann LeCun's Joint Embedding Predictive Architecture (JEPA) proposes one solution: embodied learning through physical interaction. A system learns representations via gradient descent on sensorimotor data such as from manipulating objects. This requires training hierarchical networks from scratch over extended periods, commonly on specialized hardware.

We propose an alternative learning architecture: episodic compression, that learns from narrative-formatted experiences curated based on salience rather than gradient descent.

Two Learning Architectures

Both JEPA and episodic compression are learning architectures that operate in real time:

JEPA: Learns from sensorimotor experience using gradient descent on sensory streams

Episodic Compression: Learns from narrative-formatted experience using salience-driven compression to latent variables

The key differences:

- **Input format:** Sensorimotor streams vs. narrative-formatted experiences
- **Learning mechanism:** Gradient descent vs. episodic compression
- **Implementation:** Train from scratch vs. leverage pre-trained models as cognitive substrate

Both are designed to achieve sample-efficient learning from operational experience. The question is which input format and learning mechanism proves more practical across domains.

The Episodic Compression Mechanism

Our architecture operates through three stages:

Experiences are continuous streams of interaction: conversations, observations, operational events, or simulated environments. In our proposed system, experiences fade without retention, mirroring what we experience with human memory.

Episodes are compressed summaries of experiences marked as salient. Saliency comes from surprise – when something violates expectations or predictions fail. These episodes preserve causal structure while discarding irrelevant details, naturally taking narrative form with temporal sequences and explicit cause-and-effect relationships.

Generalizations are patterns extracted when multiple episodes share structure. Our proposed system uses a threshold mechanism where around 3-5 related episodes with sufficient saliency generate a stable generalization. This explains both rapid learning from few examples and resistance to spurious patterns from single unusual events.

Episodes and generalizations serve as learned latent variables created through saliency-driven compression rather than gradient descent, enabling sample-efficient learning without parameter updates.

The architecture uses large language models as cognitive substrate. An LLM provides the computational operations (attention mechanisms for pattern-matching, reasoning capabilities for causal analysis) that the architecture orchestrates into a learning system. This isn't asking an LLM to "remember stories from training" – it's using the LLM's capabilities as infrastructure for a learning architecture.

Testing the Mechanism

Testing episodic compression faces a challenge: language models have already seen massive amounts of text during training. They know about physics, causation, social dynamics. How do you test whether a system can learn from narrative-formatted experience when it might already know the answer?

The solution: create an entirely artificial world with internally consistent but unprecedented rules (the details of the test are available on GitHub).

The Shimmer Valleys is a fictional environment with entities like Globbs (rolling, color-changing), Whisps (smoke-like, sound-responsive), and Resonators (crystalline, tone-emitting). The world operates under consistent but novel causal rules: contact transfers color through chains, sound frequencies trigger specific responses, celestial positions trigger transformations, and some changes are irreversible.

We imagined a visitor to this world, created their travelog and used an LLM to distill that into episodes. We then used another LLM on these episodes to prepare a "world model" resulting in just 9 generalizations and 4 specific high-salience episode memories – about 500 words in total, representing 65-70% compression from the original 2,000-3,000 word travelog while preserving causal structure.

With this world model we asked yet another LLM to run five novel test scenarios that would require synthesis of multiple generalizations:

- A chrome flutter seed touching a resonator
- A glowing glob entering a shade pool
- Whisps forming a chorus near a silence-sphere resonator
- A phase glob encountering a mirror surface
- Flutter seeds near micro-resonators emitting whisp-only frequencies

Results

The system successfully generated predictions for all five test scenarios (5/5) with appropriate confidence calibration (5/5). It proposed three novel mechanisms not present in the original descriptions. Confidence levels aligned with prediction basis throughout.

The system never claimed to "remember" these scenarios but instead reasoned: "Based on generalizations G1 and G5, I predict..." It referenced specific episodes rather than reconstructing from first principles. Nine generalizations enabled predictions across diverse scenarios through flexible combination.

What This Demonstrates (and Critically, What It Doesn't)

This validation shows the mechanism works at toy scale: 9 generalizations and 4 episodes can enable prediction in novel scenarios with appropriate confidence calibration.

But we need to be clear about limitations:

- **Unproven at scale:** Success with 9 generalizations doesn't guarantee success with 900. Real-world deployment requires hundreds or thousands of generalizations.
- **No adaptation validation:** The system made predictions but wasn't tested on learning from failures and updating its model accordingly.
- **Uncertain thresholds:** The proposed 3-5 episode threshold needs empirical validation across different domains.
- **Single-session only:** Long-term stability over weeks, months, or years remains completely unknown.
- **Artificial world validation:** Success in Shimmer Valleys doesn't guarantee success in messy real-world domains with noisy data and ambiguous causation.
- **Optimal parameters unknown:** Salience thresholds, consolidation frequency, and other architectural parameters require empirical tuning.

These aren't minor caveats – they're substantial open questions that determine whether this approach has practical value beyond a proof of concept.

Practical Implications (If It Scales)

If the approach scales past these limitations, it offers several operational advantages:

No Training Required: This works with existing pre-trained models. No gradient descent, no parameter updates, no server farms computing for weeks.

Immediate Deployment: Systems can be operational immediately, learning at runtime rather than requiring separate training phases.

Continuous Learning: Unlike traditional ML systems that are frozen after training, these systems improve over time during normal operation. Week 1 performance < Month 1 < Month 6.

Explainability: Episodes and generalizations are human-readable. You can inspect what the system "knows" and why it makes specific predictions.

Context-Specific Learning: Each deployment automatically learns patterns specific to its operational environment. A system deployed in one situation learns different patterns than one deployed in another situation, without manual configuration.

An interesting use case is tacit knowledge – the value comes from capturing operational knowledge that practitioners accumulate through experience but rarely document: which judge responds to which arguments, which patient symptoms predict complications, which objection patterns predict deal closure. This knowledge exists in human memory but gets lost when practitioners leave or forget. Episodic compression captures it automatically during operation.

Scaling Architecture

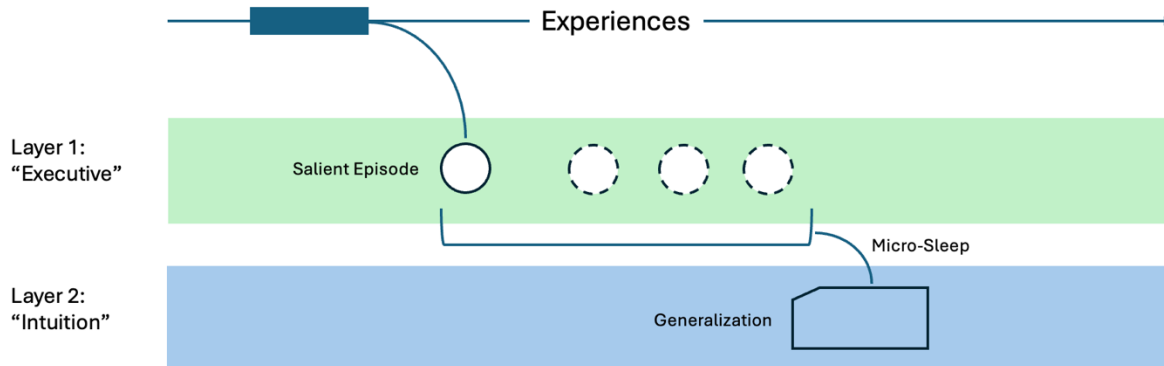
Real-world deployment requires addressing the scale problem. Our proposed solution uses a two-layer architecture:

Intuition maintains consolidated knowledge using small, efficient models (1-3B parameters) with large context windows (100K-200K tokens). It runs continuously, performing parallel pattern-matching across comprehensive knowledge bases.

Executive handles reasoning and consolidation using large, capable models with standard context windows. It engages selectively for complex queries and conducts periodic consolidation during brief "micro-sleep" cycles.

This separation exploits different computational properties: Intuition leverages transformer attention for parallel pattern-matching; Executive performs serial depth-first reasoning compressing episodes, and forms generalizations in micro-sleep cycles.

Episodic Compression



An emergent property appears: the architecture naturally operates across multiple timescales mirroring human learning. Real-time pattern-matching (seconds), episodic formation (minutes-hours), consolidation cycles (hours-days), and long-term accumulation (weeks-months-years).

Whether this architecture works at scale remains untested. That's the next research question.

Relationship to Embodied Learning

This work doesn't claim embodied learning is without merit. Physical interaction may prove optimal for tasks requiring sensorimotor grounding including robotics, manufacturing, and physical manipulation. These approaches may be complementary rather than contradictory, addressing different use cases.

JEPA learns from sensorimotor experience. Episodic compression learns from narrative-formatted experience. The Shimmer Valleys validation proves our mechanism works: the system formed generalizations about novel causal rules, then combined those generalizations to predict outcomes in scenarios it had never encountered.

In practical deployment, narrative-formatted experiences come from operational activity:

- Direct observations the system makes during operation
- Interactions with users (conversations, queries, feedback)
- Operational outcomes and their contexts
- Case studies describing domain-specific events

The system applies episodic compression regardless of source: salient experiences compress to episodes, episodes accumulate, pattern detection extracts commonalities, generalizations form when multiple episodes share structure.

This enables continuous learning during deployment. System performance improves over time – week 1 < month 1 < month 6 – from accumulated episodes and generalizations specific to the operational context, not from parameter updates.

Whether learning from narrative-formatted experience proves optimal for a given domain, or whether sensorimotor experience provides advantages, remains an empirical question. Our contribution is demonstrating that episodic compression enables sample-efficient learning from narrative-formatted input using existing technology.

The Practical Question

JEPA requires training hierarchical networks from scratch over extended periods on specialized hardware. Episodic compression can be deployed today using off-the-shelf LLMs as cognitive substrate.

The architecture described here – episodic compression to latent variables – offers a practical path to sample-efficient learning using tools we already have. Systems that learn continuously while deployed, adapt to specific operational contexts automatically, and accumulate expertise without expensive retraining cycles.

Whether our proposed learning system proves optimal or complementary to other approaches, episodic compression provides a viable path forward with immediate practical value if it scales past the toy validation we've demonstrated so far.

The initial validation suggests the mechanism is sound. Everything else requires further research.

This article is based on research from Mossrake Group, LLC on episodic compression and narrative-based learning. The full technical papers that explore implementation details, scaling architectures, and multi-timescale learning mechanisms can be found at <https://github.com/mossrake/learning-system>.

© 2025 Mossrake Group, LLC

Version 1.0