

Wordscillation: Learned Energy Ranking of Bounded Coherent Hypotheses in Compact Structured Memory

Jaharri Gassaway (aka DumbButt)
Inquiry Labs

September 1, 2026

Abstract

When a compact memory is noisy, should a reasoner immediately decode each part into one symbol, or preserve several complete interpretations until their relations can be judged together? We study this question in a controlled synthetic setting. Each memory packet binds four roles—source, relation, context, and target—under a fixed budget of 65,536 real scalars. A query requires traversing a four-hop path through a four-branch world. Rather than committing to independently decoded roles, our system retrieves a small packet shortlist, constructs a width-8 beam of coherent whole-path hypotheses, and ranks those hypotheses with an identity-free learned energy function. On a registered within-family evaluation, a 1,153-parameter multilayer perceptron improves dense moderate-corruption accuracy from 82.188% for a frozen structural score to 94.896% (paired gain 12.708 percentage points; 95% seed-cluster interval 9.688–15.521), close to a 95.625% retained-path oracle. A separately registered fixed, non-oracle top-16 retriever matches the 95.729% oracle-shortlist ceiling. The result does not transfer without adaptation to a qualified 32-chain $\times$ 2-branch family: accuracy is 28.333%, near both the frozen score (25.625%) and shuffled-label control (25.208%), despite a 78.125% oracle ceiling. Post-registered exploratory tests sharpen this boundary. A family-invariant scorer improves qualified depth and memory-load extrapolations but fails a prespecified general-transfer gate because branch transfer remains negative and the joint extreme is unqualified. A parameter-matched message-passing graph also fails to repair the qualified branch case, reaching 23.333% versus 61.667% for the coherent-path beam. We therefore report a narrow positive result and a decisive boundary. Learned energy can resolve structured ambiguity within the tested family; it is not yet a general reasoning rule, a language model, or evidence for an oscillatory mechanism. The broader Wordscillation intuition survives only in a refined form: uncertainty should remain structured long enough for complete hypotheses to compete before commitment.

1 Introduction

Many model papers begin from an architecture and ask how far it scales. We began from a failure.

The original Wordscillation idea was that uncertain input should not collapse immediately to one meaning. A partial observation such as *melo* might support several continuations, and later context should be allowed to strengthen or weaken them. Early toy experiments represented those possibilities as waves. That metaphor was useful because it insisted on continuous uncertainty, but the experiments did not establish that oscillation, phase synchronization, or wave coherence caused better reasoning.

The project became more precise when we moved from words to role-bound relational memories. A packet now represents a structured event with source, relation, context, and target roles.

Compression makes many such packets cheap to store, but corruption makes their contents ambiguous. Hard cleanup—decoding each role to its nearest symbol before reasoning—often destroys information that the other roles could have disambiguated. Independent soft marginals preserve uncertainty but can combine locally plausible fragments into no globally coherent explanation. The key question therefore became:

Can a compact reasoner preserve a bounded set of complete hypotheses and learn which whole explanation is supported?

This paper answers that question in one synthetic family and tests the answer outside it. Our system has four deliberately separated components: compact role-bound packet memory, query-conditioned non-oracle retrieval, a bounded beam of complete four-hop paths, and a small identity-free energy scorer. The scorer never sees entity identities. It receives 70 structural features describing unary role evidence, cross-packet transitions, and ambiguity margins, and assigns one scalar energy to each retained path.

The main within-family result is strong but narrow. Under moderate corruption, the learned scorer reaches 94.896% accuracy versus 82.188% for the frozen structural score. The gain is paired across ten registered data-generation seeds and persists across three frozen model initializations. A shuffled-label model collapses to 13.125%, while a 71-parameter linear scorer is already strong, showing that the useful signal is learned structural weighting rather than capacity alone. Separately, a fixed top-16 retriever matches the oracle shortlist on another registered holdout.

The most important result may be the one that fails. On a qualified structural transfer from 16 chains $\times$ 4 branches to 32 chains $\times$ 2 branches, the correct path remains available often enough to create 52.5 points of oracle headroom, yet the learned scorer recovers only 2.708 points over the frozen score. The learned energy is therefore family-specific under this change.

We then ran two explicitly exploratory, post-registered mechanism checks. P1-X2 replaces the original 70-feature interface with a 61-feature family-invariant interface and trains once on a mixed grid of depths, branch factors, and memory loads. It transfers strongly across unseen depths and loads but still fails its general-transfer gate because branch-factor performance is negative, one joint cell is unqualified, and the worst qualified cell loses five points. P1-X3 asks whether a comparably sized packet graph can repair that branch weakness. It cannot: on the qualified branch cell, the graph underperforms both coherent-path ranking and its own no-message ablation. These later results do not revise the registered claim; they narrow the mechanism and its present boundary.

None of the ingredients alone is presented as novel: role-filler binding, compact vector-symbolic memory, beam search, energy-based ranking, and multi-hop path inference all have substantial prior art. The paper instead isolates a specific mechanistic claim and experimental configuration: under corruption in equal-budget compact role-bound memory, candidate survival should be separated from candidate selection by retaining a small beam of complete paths and ranking them late with an identity-free structural energy. Put differently, uncertainty over isolated role fillers is not equivalent to uncertainty over complete relational explanations.

Our contributions are:

1. A controlled equal-budget task for studying ambiguity in compact role-bound memory.
2. A bounded coherent-hypothesis mechanism that separates candidate survival from candidate selection.
3. Registered evidence that a very small identity-free energy scorer improves path ranking within one family.
4. Separately registered evidence that fixed non-oracle retrieval can match an oracle shortlist on the same task family.

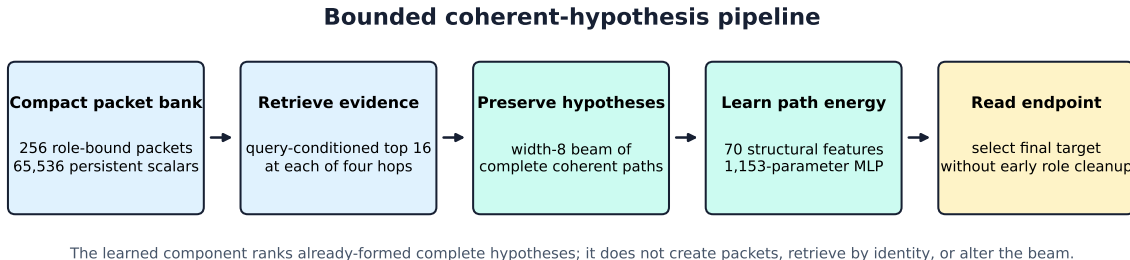


Figure 1: Figure 1. The tested pipeline. Compact role-bound packets are retrieved into a top-16 shortlist. A width-8 beam preserves complete path hypotheses. An identity-free scorer ranks the beam and the selected path supplies the endpoint answer.

5. A qualified negative transfer result that defines what the learned score does not generalize to.
6. Post-registered exploratory evidence separating successful depth/load transfer from unresolved branch transfer and rejecting one parameter-matched graph-message repair.

We do not claim a new binding algebra, a first bounded-hypothesis reasoner, general multi-hop reasoning, state-of-the-art efficiency, natural-language understanding, adaptive computation, or an oscillatory advantage. The contribution is a mechanistic study of late commitment under compact structured uncertainty.

2 Related Work

2.1 Role-filler binding and vector-symbolic memory

Tensor Product Representations (TPRs) formalize structured representation by binding fillers to roles and superposing the resulting products [1]. Holographic Reduced Representations (HRRs) compress related structure into fixed-width vectors using circular convolution, accepting noisy reconstruction and cleanup as part of the tradeoff [2]. Modern surveys place TPR, HRR, MAP, and related systems inside the broader hyperdimensional-computing and vector-symbolic-architecture family [3, 4]. Differentiable HRR work has revisited stability and learnability [5], while Hrrformer recasts attention-like query/key/value processing through HRR operations for long sequences [6]. Alam et al. also develop a differentiable Walsh–Hadamard-derived binding operation [7].

Our study does not propose a new binding operator. It uses an equal-budget comparison to expose a familiar compression tradeoff: compact packets store more experiences, but noisy role evidence becomes harder to resolve. Our focus is the inference procedure after a compact packet has become ambiguous.

2.2 Energy-based structured selection

Energy-based learning assigns scalar compatibility to candidate configurations and selects low-energy or high-score states [8]. Our $70 \rightarrow 16 \rightarrow 1$ network is a small supervised energy scorer in this established sense. What is specific here is its interface: it receives no entity identity and ranks a bounded set of already coherent multi-packet paths. This lets us ask whether local evidence can be recombined into a useful global judgment without hiding the mechanism inside a large encoder.

2.3 Multi-hop and memory reasoning

Differentiable proving, DeepPath, and later knowledge-graph work learn to traverse or score relational paths [9–11]. These systems operate in richer and more realistic settings than ours. Our task is intentionally smaller because the goal is causal isolation: hold the memory budget, path depth, beam width, and feature interface fixed, then determine whether the ranking failure is recoverable.

Recent sequence models also motivate explicit study of memory. Transformers model direct dependencies through attention [12]; Mamba introduces input-dependent selective state-space dynamics to address content-based reasoning while retaining linear sequence scaling [13]; Titans frames attention and neural memory as short- and long-term components [14]. Those models are not direct baselines here because this paper does not test end-to-end sequence modeling. They do, however, demonstrate a productive style of research: identify one bottleneck, define one mechanism, and test that mechanism against clear alternatives. We follow that structure at a much smaller scale.

2.4 Why the project name is not the present claim

Rotary and oscillatory mechanisms already have substantial prior art, including RoPE [15], RotRNN [16], rotational recurrent units, unitary recurrence, Selective Synchronization Attention [17], and Attention as Frustrated Synchronization [18]. The last two are recent 2026 preprints and their empirical claims should be treated as preliminary. The experiments reported here contain no learned frequency, phase coupling, or synchronization process. “Wordscillation” names the research lineage and the original intuition that uncertainty should evolve before commitment. The implemented paper is a structured-memory and energy-ranking study.

3 Problem Setting

3.1 Synthetic relational worlds

Each base world contains 16 independent chains. Every query follows one intended path of depth four. At each layer, four branches compete, producing distractor packets that are locally plausible but globally incompatible with the intended chain. A packet is a tuple

$$p_i = (s_i, r_i, c_i, t_i),$$

where s , r , c , and t denote source, relation, context, and target. The full world contains 256 packets. A query provides an initial source together with the required relation and context sequence; the answer is the endpoint of the supported four-hop path.

Train, development, and registered test worlds use disjoint data-generation seeds and newly sampled codebooks. The learned scorer therefore cannot solve the task by memorizing entity vectors from training.

3.2 Equal persistent-memory budget

Persistent packet storage is fixed at 65,536 real scalars. A full TPR-style packet consumes 1,024 scalars, so only 64 packets, or four complete chains, fit. A compact packet consumes 256 scalars, so all 256 packets, or 16 chains, fit. The paper’s primary learned-ranking result uses the dense compact representation, which concatenates four 64-dimensional role blocks. A compressed MAP variant superposes four role-masked 256-dimensional fillers and is treated as secondary.

This accounting concerns persistent storage only. Temporary workspace, trainable parameters, runtime, FLOPs, and wall-clock latency are separate resources and are not silently folded into the same number.

3.3 Corruption

Packet vectors are corrupted while preserving norm and controlling cosine similarity to the clean packet. We report clean and moderate-corruption conditions, with moderate corruption serving as the primary regime. The task is trivial when clean: the interesting question is whether structure can recover information that local cleanup loses under noise.

3.4 What counts as success

The intended path must first survive retrieval and beam construction. Only then can a scorer select it. We therefore separate:

- **packet-path recall:** whether retrieval exposes every packet on the intended path;
- **gold-path beam retention:** whether the complete intended path remains in the width-8 beam;
- **gold-path selection:** whether a scorer ranks that path first;
- **answer accuracy:** whether the selected endpoint is correct.

This distinction prevents a high oracle score from being mistaken for a working model, and it prevents a low model score from being blamed on ranking when the candidate was already removed.

4 Method

4.1 Local role evidence

For every packet, the model compares each noisy role block against its clean role codebook. Cosine similarities are converted to role likelihoods with a temperature-scaled softmax. This yields distributions over candidate fillers for source, relation, context, and target without an immediate hard cleanup decision.

4.2 Coherent path beam

At layer ℓ , a partial path receives unary evidence from the expected relation and context. The first layer also receives evidence for the query source. Adjacent packets are connected by a transition score measuring compatibility between the previous packet’s target distribution and the next packet’s source distribution:

$$T(p_\ell, p_{\ell+1}) = \log \left(q_\ell^{(t)} \cdot q_{\ell+1}^{(s)} + \epsilon \right).$$

The frozen path score adds unary and transition terms across four layers. Beam search retains the eight highest-scoring complete hypotheses. Unlike independent marginals, each beam element is a coherent packet sequence with one endpoint.

4.3 Identity-free path features

For each retained path, we extract 70 real-valued features. They include per-layer unary relation and context terms, first-layer source evidence, adjacent transition terms, aggregates over these components, selected-versus-alternative margins, and distribution summaries such as maximum

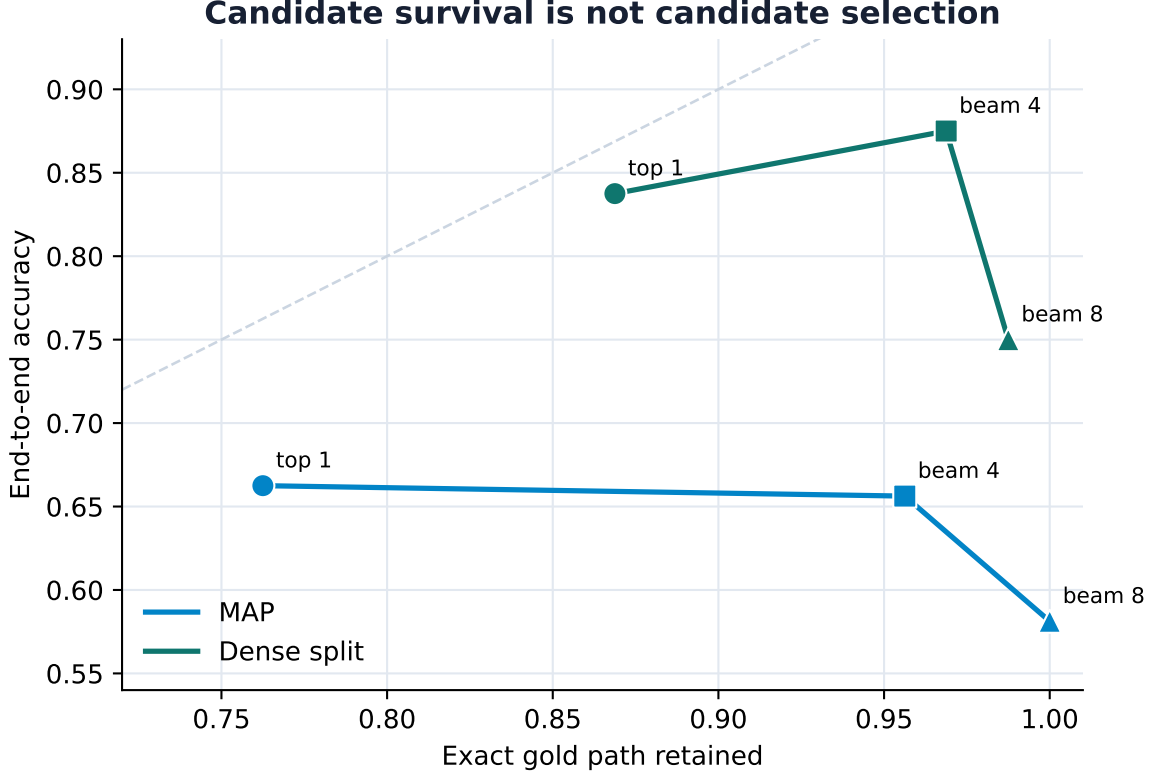


Figure 2: Figure 2. Candidate survival is necessary but not sufficient. W2K-E showed that the correct path could remain in the beam while the frozen score still selected the wrong path. W2K-F targets this gap directly.

probability, top-two margin, and entropy for each role. Entity IDs and raw code vectors are excluded.

This restriction is central. The scorer can learn how evidence should be weighted, but it cannot memorize that a particular named entity tends to be correct.

4.4 Learned energy

The primary scorer is

$$E_{\theta}(h) = w_2^{\top} \tanh(W_1 f(h) + b_1) + b_2,$$

with 70 inputs, 16 hidden units, one output, and 1,153 trainable parameters. A 71-parameter linear scorer tests whether nonlinear capacity is necessary. Training minimizes listwise cross-entropy over the eight retained hypotheses. Features are normalized using the training moderate-corruption split only.

Training uses full-batch AdamW for exactly 300 epochs, learning rate 0.01, and weight decay 10^{-4} . There is no early stopping or checkpoint selection. We train on 20 data seeds, inspect five development seeds, and evaluate once on ten registered test seeds. Three frozen initialization seeds are used for each learned model.

4.5 Retrieval

The initial W2K-F ranking study evaluates a controlled oracle packet shortlist so that ranking can be isolated. P1-G1 then removes that oracle. At each hop, the fixed retriever scores every packet by the sum of log matches for current source state, expected relation, and expected context, retains the top 16, and advances the current state using weighted target evidence. The reasoner and its learned scorer remain frozen.

We compare fixed retrieval with a learned retriever, an oracle shortlist, and an equal-size random shortlist. The learned retriever is a control, not the headline: the experiment asks whether oracle retrieval can be removed without losing the ranking result.

5 Experimental Protocol

5.1 Baselines and controls

- **Frozen top path:** the highest fixed structural score; primary ranking baseline.
- **Soft local:** inherited independent local-inference control.
- **Linear scorer:** learned identity-free $70 \rightarrow 1$ score.
- **MLP scorer:** learned identity-free $70 \rightarrow 16 \rightarrow 1$ score.
- **Shuffled-label MLP:** identical capacity trained on deterministically wrong labels.
- **Gold-path oracle:** chooses the intended path when retained; an upper bound after beam construction.
- **Fixed non-oracle retriever:** primary registered retrieval method.
- **Random top-16 retriever:** equal-shortlist retrieval floor.

5.2 Statistical unit and uncertainty

The independent uncertainty unit is the data-generation seed. The registered W2K-F set contains ten seed clusters, two worlds and 32 queries per seed, for 320 queries. Outcomes are averaged across the three model initializations and then within each data seed. We report two-sided 95% percentile intervals from 50,000 seed-cluster bootstrap replicates and exact two-sided sign-flip tests over the ten paired seed effects.

These inferential analyses are post-hoc and descriptive. They were specified after the registered aggregate result was known and are not multiplicity-adjusted. The registered aggregate itself remains frozen and byte-verified.

5.3 Structural-transfer stress test

The first transfer family was uninterpretable because the width-8 beam often removed the intended path. We did not relabel this as a model failure. Instead, a second transfer family, G2-B, changes the structure from 16 chains $\times$ 4 branches to 32 chains $\times$ 2 branches while requiring complete fixed-retrieval packet-path recall and adequate gold-path beam retention. Only after this qualification do we judge transfer.

5.4 Reproducibility and frozen evidence

All protocols, outputs, and final audit inputs are content-hashed. The final audit verified eight frozen input artifacts byte for byte and reconstructed the consumed registered split only to obtain paired outcomes. It allocated no new experimental seeds and created no replacement task result.

System	Accuracy	95% seed-cluster interval
Frozen top path	82.188%	79.375–85.000%
Soft local	89.063%	87.500–90.938%
W2K-F MLP	94.896%	94.063–95.833%
Gold-path oracle	95.625%	94.688–96.563%
Shuffled-label MLP	13.125%	10.417–16.250%

The final audit result SHA-256 is 9f1dfcd609dc372510e6d6f38079b2eecb1bf3ebf8d6558d92d8f84ceaeb6a90.

5.5 Post-registered exploratory extensions

P1-X2 asks whether one identity-free scorer can transfer across a factorial structural grid. One pooled model is trained on depths 3–5, branch factors 2–4, and memory loads 128–256, with four complete compositions removed. Evaluation includes those held-out compositions, depth 6, branch factor 6, memory load 512, and their joint extreme. The fixed top-16 retriever, width-8 beam, moderate-corruption primary regime, and per-cell qualification rules remain unchanged. The X2 MLP has 61 inputs, 18 hidden units, and 1,135 parameters; a 62-parameter linear scorer and a shuffled-label MLP are controls.

P1-X3 is a post-X2 mechanism test designed after branch weakness was observed. It replaces explicit path enumeration with a layered graph over the same top-16 packet shortlist. Each packet has 19 identity-free local features. Adjacent-layer target-to-source compatibility weights messages in a shared 19→24 local encoder, 24→24 message transform, and 24→1 final-packet readout, totaling 1,105 parameters. A parameter-matched no-message ablation and wrong-label message graph test whether interaction and supervision matter. Because X2 and X3 were designed after earlier outcomes were known, they are reported as exploratory development evidence, not registered confirmation.

6 Results

6.1 Learned ranking closes most of the within-family gap

Table 1 reports the registered dense moderate-corruption result.

The MLP improves over the frozen score by 12.708 percentage points (95% interval 9.688–15.521; exact sign-flip $p = 0.001953$). It improves over soft local inference by 5.833 points (4.583–7.083) and reduces the observed frozen-score error by 71.345% (64.151–77.222%). Its 0.729-point deficit to the retained-path oracle has an interval reaching zero (−1.563 to 0.000).

The shuffled-label result matters as much as the parameter count: the same architecture trained on wrong supervision does not inherit the gain. The strong linear control further indicates that the task does not require a deep nonlinear model. The main mechanism is learned weighting of structural evidence.

6.2 Fixed retrieval removes the oracle dependency

On a separate registered holdout, fixed, learned, and oracle retrieval all reach 95.729% mean accuracy and 100% complete packet-path recall. Random top-16 retrieval reaches 5.937% and 0% recall. The same result holds under all three frozen reasoner initializations.

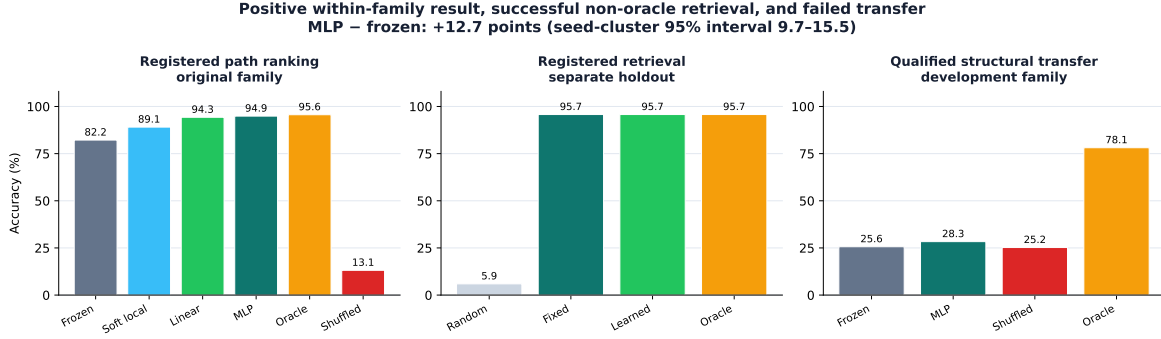


Figure 3: Figure 3. Left: learned ranking succeeds within family. Center: fixed non-oracle retrieval matches the oracle shortlist. Right: the learned ranking advantage does not transfer to the qualified G2-B family.

This establishes that the within-family result does not require an oracle packet shortlist. It does not establish that learned retrieval is superior; the fixed rule is equally effective on this family.

6.3 Structural transfer fails after qualification

G2-B achieves 100% packet-path recall and 76.875% gold-path beam retention. The oracle therefore reaches 78.125%, exposing 52.5 points of headroom over the frozen score. Nevertheless, the frozen W2K-F MLP reaches only 28.333%, compared with 25.625% for the frozen structural score and 25.208% for the shuffled-label control.

The within-family scorer has learned a useful regularity, but not one invariant to this tested structural change. This is not a footnote: it is the current boundary of the contribution.

6.4 Exploratory structural sweep narrows the transfer boundary

P1-X2 fails its conjunctive general-transfer gate, but the failure is structured rather than uniform. Ten of eleven moderate-corruption cells qualify. Across those qualified cells, the invariant MLP improves over frozen ranking by 17.333 points on average and beats its shuffled-label control by 72 points. Mean single-axis gains are +37.5 points for unseen depth and +27.5 for unseen memory load, but −3.333 for unseen branch factor. The worst qualified-cell gain is −5 points. The joint depth-6 × branch-6 × load-512 cell is unqualified because retrieval and beam survival collapse.

These results show that the original one-family failure is not a universal inability to extrapolate. The learned structural weighting can survive substantial changes in depth and stored-memory load. It has not become a general inference rule because increased branching remains an unresolved failure and the joint extreme does not expose a qualified ranking problem.

6.5 Message passing does not repair qualified branch transfer

P1-X3 compares the 1,105-parameter graph with the 1,135-parameter X2 beam MLP. Only one of the two prespecified branch cells qualifies under the unchanged X2 contract. On that qualified depth-5, branch-6, load-128 cell, fixed retrieval is complete, beam retention is 95%, and both the retained-path and graph endpoint oracles provide adequate headroom. Mean accuracies are:

The graph therefore trails the beam by 38.333 points and its own no-message ablation by 5 points. Across all nine qualified moderate-corruption holdouts, the message graph reaches 79.259%,

System	Accuracy
X2 invariant beam MLP	61.667%
Frozen structural ranking	35.000%
Graph without messages	28.333%
Graph with messages	23.333%
Shuffled-label message graph	8.333%

Resource	Budget
Persistent packet memory	65,536 real scalars
Stored packets / chains	256 / 16
Coherent workspace upper bound	11,101 real scalars
Eight paths $\times$ 70 scorer features	560 real scalars
MLP parameters	1,153
Linear parameters	71
Retrieval shortlist	16 packets per hop
Hypothesis beam	8 complete paths

the beam reaches 88.519%, and the no-message graph reaches 60.741%; the graph beats the beam on zero of nine cells. Messages carry useful information in aggregate, but endpoint-scored message passing does not replace coherent whole-path competition in this experiment.

6.6 Resource accounting

These numbers demonstrate bounded storage and model size. They are not an end-to-end latency, FLOP, energy-use, or hardware-efficiency comparison.

7 What the Experiment Changed

The original intuition was “let possibilities behave like waves until context makes them settle.” The experiments changed almost every noun in that sentence while preserving its core warning.

First, literal wave dynamics did not earn the central role. Second, hard nearest-symbol cleanup failed because it committed before all the evidence could interact. Third, independent soft distributions also failed because uncertainty over isolated roles is not the same thing as uncertainty over complete explanations. What survived was more specific:

Keep a small number of coherent interpretations alive, let shared relational evidence change their relative support, and commit only after comparing the complete hypotheses.

This is less poetic than the starting idea, but more testable. “Amplitude” becomes hypothesis weight. “Interference” becomes competition under a learned energy. “Settling” becomes a bounded selection process. These are analogies, not claims about physical oscillation.

The failed transfer result is equally constructive. It says that the current 70-feature interface contains family-specific assumptions. A future model should not simply scale this scorer and call the problem solved. It should identify which evidence transformations remain stable when branch factor, memory load, path length, and relational grammar change.

The exploratory extensions refine that statement. A smaller invariant interface is already stable across unseen depth and memory load, so the failure is not simply “all structure changes break the scorer.” Increased branching is the sharper boundary. The X3 result also rejects one tempting shortcut: diffusing transition messages over all shortlisted packets and reading out an endpoint

does not preserve the advantage of competing complete explanations. Message passing may still be useful, but in this setting it helps relative to a local-only graph while remaining weaker than explicit whole-hypothesis ranking.

8 Limitations and Claim Boundary

This study has substantial limitations.

Synthetic task. The worlds have known roles, fixed depth, controlled distractors, and generated codebooks. No conclusion about natural-language understanding or generation follows.

Hand-specified interface. The event schema, role boundaries, beam procedure, and 70 features are designed rather than learned end to end.

Family-specific learning. The principal learned advantage fails under the qualified G2-B structural change.

Exploratory extensions. X2 and X3 were designed after earlier transfer outcomes were observed. Their frozen development runs diagnose the boundary but are not independent registered confirmation.

Branch generalization. The invariant X2 scorer transfers across depth and memory load but not increased branching. The endpoint-scored X3 graph does not repair the qualified branch case.

Oracle remains in one diagnostic. W2K-F uses oracle shortlists to isolate ranking, although P1-G1 separately demonstrates a fixed non-oracle replacement on a registered holdout.

Bounded beam. Beam width eight is part of the tested mechanism. Larger or more difficult worlds may discard the correct path before scoring.

Post-hoc uncertainty analysis. Confidence intervals and sign-flip tests summarize frozen outcomes but were not preregistered confirmatory analyses.

No efficiency claim. Scalar budgets and parameter counts do not substitute for FLOP, latency, memory-bandwidth, or energy measurements.

No oscillatory mechanism. The paper provides no evidence that frequency, phase, synchronization, Fourier dynamics, or rotational recurrence causes the observed ranking gain.

Accordingly, the supported claim is narrow: within one controlled synthetic family, a small identity-free learned energy ranks bounded coherent path hypotheses substantially better than a frozen structural score after relevant evidence has been exposed. The registered score does not transfer under the tested family change. Exploratory evidence shows that an invariant redesign can transfer across depth and memory load, but neither that redesign nor a parameter-matched message graph resolves qualified branch-factor generalization.

9 Conclusion

Compact memory creates a temptation to decode early. Our experiments suggest that this can be exactly the wrong move when evidence is distributed across a relation. The correct interpretation may be weak in every isolated role yet strong as a complete path.

Wordscillation began as an attempt to keep meanings fluid. Paper 1 does not validate waves, attention replacement, or a new language architecture. It does establish a smaller principle in a setting where we can see the mechanism: preserve a bounded set of coherent hypotheses, then learn how the whole evidence should be ranked.

That principle worked within the family in which it was learned, nearly closing the retained-path oracle gap. It also failed to transfer when the structure changed. Later exploratory tests found that depth and memory-load extrapolation are attainable, while increased branching remains the

decisive weakness. A small message-passing graph did not repair it. A solid conclusion must contain all three facts. The next research question is not whether to hide the failure with a larger model. It is how to preserve diverse complete explanations as branching grows while retaining the transparency, bounded workspace, and late-commitment behavior demonstrated here.

References

- [1] P. Smolensky. “Tensor product variable binding and the representation of symbolic structures in connectionist systems.” *Artificial Intelligence*, 1990. [https://doi.org/10.1016/0004-3702\(90\)90007-M](https://doi.org/10.1016/0004-3702(90)90007-M)
- [2] T. A. Plate. “Holographic Reduced Representations.” *IEEE Transactions on Neural Networks*, 1995. <https://doi.org/10.1109/72.377968>
- [3] D. Kleyko, D. A. Rachkovskij, E. Osipov, and A. Rahimi. “A Survey on Hyperdimensional Computing aka Vector Symbolic Architectures, Part I.” *ACM Computing Surveys*, 2022. <https://arxiv.org/abs/2111.06077>
- [4] D. Kleyko, D. A. Rachkovskij, E. Osipov, and A. Rahimi. “A Survey on Hyperdimensional Computing aka Vector Symbolic Architectures, Part II.” *ACM Computing Surveys*, 2023. <https://arxiv.org/abs/2112.15424>
- [5] A. Ganesan et al. “Learning with Holographic Reduced Representations.” *NeurIPS*, 2021. <https://arxiv.org/abs/2109.02157>
- [6] M. M. Alam, E. Raff, S. Biderman, T. Oates, and J. Holt. “Recasting Self-Attention with Holographic Reduced Representations.” *ICML*, 2023. <https://arxiv.org/abs/2305.19534>
- [7] M. M. Alam et al. “A Walsh Hadamard Derived Linear Vector Symbolic Architecture.” *NeurIPS*, 2024. https://proceedings.neurips.cc/paper_files/paper/2024/hash/0525fa17a8dbea687359116d01732e12-Abstract-Conference.html
- [8] Y. LeCun, S. Chopra, R. Hadsell, M. Ranzato, and F. Huang. “A Tutorial on Energy-Based Learning.” 2006. <https://yann.lecun.org/exdb/publis/pdf/lecun-06.pdf>
- [9] T. Rocktäschel and S. Riedel. “End-to-end Differentiable Proving.” *NeurIPS*, 2017. <https://proceedings.neurips.cc/paper/2017/hash/b2ab001909a8a6f04b51920306046ce5-Abstract.html>
- [10] W. Xiong, T. Hoang, and W. Y. Wang. “DeepPath: A Reinforcement Learning Method for Knowledge Graph Reasoning.” *EMNLP*, 2017. <https://aclanthology.org/D17-1060/>
- [11] X. V. Lin, R. Socher, and C. Xiong. “Multi-Hop Knowledge Graph Reasoning with Reward Shaping.” *EMNLP*, 2018. <https://aclanthology.org/D18-1362/>
- [12] A. Vaswani et al. “Attention Is All You Need.” *NeurIPS*, 2017. https://proceedings.neurips.cc/paper_files/paper/2017/hash/3f5ee243547dee91fbd053c1c4a845aa-Abstract.html
- [13] A. Gu and T. Dao. “Mamba: Linear-Time Sequence Modeling with Selective State Spaces.” arXiv:2312.00752, 2023. <https://arxiv.org/abs/2312.00752>

Gate	Observation	Consequence
W2K-C	Hard cleanup damaged compact representations	Preserve graded evidence
W2K-D	Independent soft marginals diffused ambiguity	Maintain joint interpretations
W2K-E	Gold path survived, frozen score misranked it	Learn a path-level score
W2K-F	Learned identity-free energy succeeded within family	Freeze scorer and test retrieval
P1-G1	Fixed top-16 retrieval matched oracle	Remove oracle dependency within family
G2	Beam qualification failed	Do not interpret scorer performance
G2-B	Task qualified; learned transfer failed	Bound the claim to one family
P1-G3	Frozen audit reproduced evidence and uncertainty	Close Paper 1 research gates
P1-X2	Invariant scorer transferred across depth/load but not branch factor	Localize the structural boundary
P1-X3	Message graph failed to repair the qualified branch case	Keep complete hypotheses as the current core mechanism

- [14] A. Behrouz, P. Zhong, and V. Mirrokni. “Titans: Learning to Memorize at Test Time.” arXiv:2501.00663, 2024. <https://arxiv.org/abs/2501.00663>
- [15] J. Su et al. “RoFormer: Enhanced Transformer with Rotary Position Embedding.” arXiv:2104.09864, 2021. <https://arxiv.org/abs/2104.09864>
- [16] K. Biegun, R. Dolga, J. Cunningham, and D. Barber. “RotRNN: Modelling Long Sequences with Rotations.” arXiv:2407.07239, 2024; revised 2026. <https://arxiv.org/abs/2407.07239>
- [17] H. Hays. “Selective Synchronization Attention.” arXiv:2602.14445, 2026. <https://arxiv.org/abs/2602.14445>
- [18] J. Nunley. “Attention as Frustrated Synchronization.” arXiv:2606.18694, 2026. <https://arxiv.org/abs/2606.18694>

A Experiment Lineage

The final mechanism was not selected in one jump.

B Frozen Splits and Training Constants

- Training data seeds: 87101–87120.
- Development data seeds: 87201–87205.
- Registered test data seeds: 87301–87310.
- Model initialization seeds: 87401–87403.
- Beam width: 8.
- Feature dimension: 70.
- MLP hidden dimension: 16.
- Epochs: 300.
- Learning rate: 0.01.

- AdamW weight decay: 0.0001.
- Primary corruption regime: moderate.
- No early stopping or checkpoint selection.

C Intended Use of the Name “Wordscillation”

The name records the project’s origin: uncertainty was first imagined as several wave-like possibilities that evolve until evidence supports one. In the present paper, no literal oscillator is used. The closest operational equivalent is a distribution over bounded coherent hypotheses and a learned energy that changes their ranking. Future work may compare dense, complex, phase/amplitude, and other geometries, but those comparisons are outside this paper.