

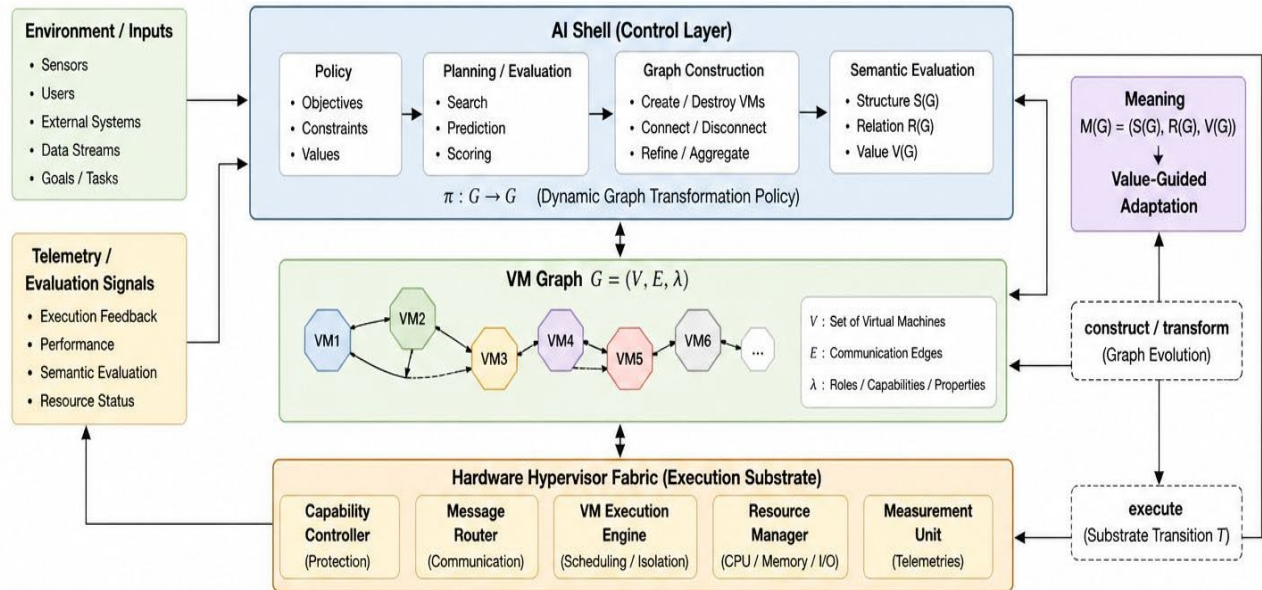
RESEARCH MEMO: QUANTUM HYPER-AETHER (1)

Hirofumi Inomata

2026

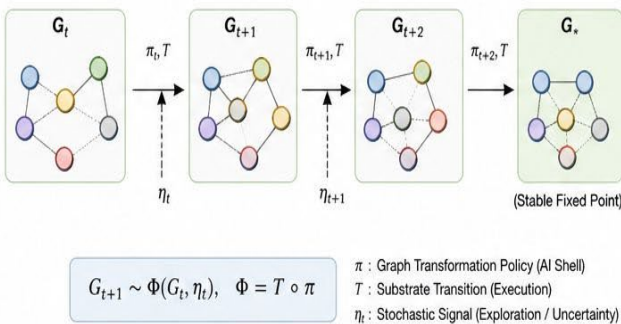
Several thin, white, parallel lines of varying lengths and slopes are positioned on the right side of the slide, extending from the top right towards the bottom left.

Figure 1. Overall Architecture of Hyper-Aether



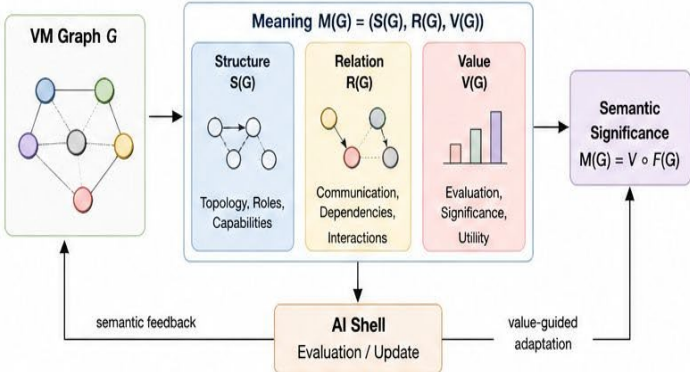
Hyper-Aether integrates AI-driven graph construction with hardware-level virtualization. Virtual Machines (VMs) are the primary semantic computational units.

Figure 2. Computation as Stochastic Graph Dynamics



At each step, the AI Shell applies a graph transformation policy π , and the hardware substrate executes the resulting graph via transition T . Stochastic signals η_t enable exploration and adaptation. Computation is the stochastic evolution of VM graphs.

Figure 3. Emergence of Meaning from VM Graphs



Meaning emerges from the joint organization of structure, relation, and value. Semantic evaluation feeds back into the AI Shell, guiding the construction and transformation of VM graphs.

概要 本稿では、仮想マシンによる実行、ハードウェアレベルの仮想化、および能力ベースの保護を動的論理システムの形式理論と統合する、AIネイティブなコンピュータアーキテクチャであるHyper-Aetherを提案する。我々は、人工知能をゲーデル不完全性によって制約された形式システムの時間発展シーケンスとしてモデル化し、計算が構造化された意味グラフの進化として再解釈できることを示す。Hyper-Aetherはこのモデルを物理的に実現する。仮想マシンはグラフ構造の計算オブジェクトを形成し、AI駆動の制御層がこれらのグラフを動的に構築および変換する。我々は、実行をグラフ空間上の確率的動的システムとして形式化し、構造-関係-値のトリプルに基づく構成的意味論を導入し、アーキテクチャを圏論的モデルおよび計算モデルに結びつける。本研究は、計算を値誘導型構造形成として再定義する。

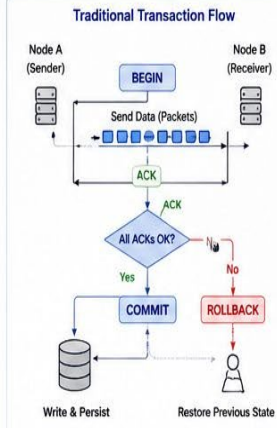
Transaction in Communication: Traditional vs Hyper-Aether (Semantic Transaction)

From bit-level consistency to semantic consistency

Traditional Transaction (Bit-Consistent World)

Goal: Ensure bit-exact delivery and state consistency

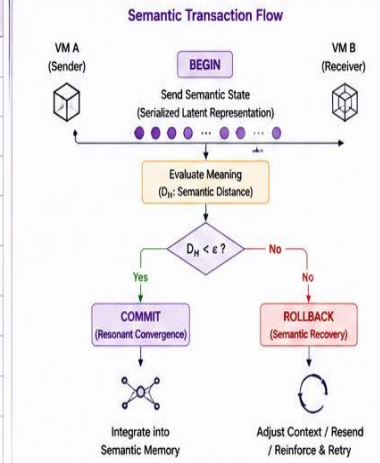
Aspect	Traditional Approach
Unit of Transaction	Packet / Record / Message
Consistency	Bit-level consistency
Success Condition	ACK received / All packets delivered
Rollback Trigger	Timeout / Error / NACK
Rollback Method	Retransmit / Restore previous state
Commit Condition	All data written and acknowledged
Isolation	Lock / Sequence / Version control
Durability	Persist to stable storage
Focus	Data correctness
Example	Database transaction, TCP, 2PC



Hyper-Aether Semantic Transaction (Meaning-Consistent World)

Goal: Ensure semantic alignment and cognitive state coherence

Aspect	Hyper-Aether Approach
Unit of Transaction	Semantic State (Meaning Unit)
Consistency	Semantic consistency
Success Condition	Semantic convergence (Meaning alignment)
Rollback Trigger	Semantic divergence ($D_H > \epsilon$)
Rollback Method	Context resend / Attention adjustment Memory correction / Reinforcement
Commit Condition	Semantic distance $D_H < \epsilon$ (Resonant convergence)
Isolation	Context boundary / Attention scope
Durability	Store in semantic memory (graph/vector)
Focus	Meaning correctness (Understanding match)
Example	Distributed cognition, Semantic OS, VM communication



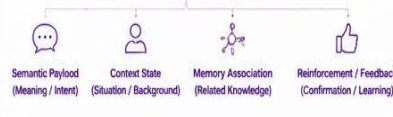
Detailed Comparison

Commit is decided by bit-level delivery, the transaction is still "successful".

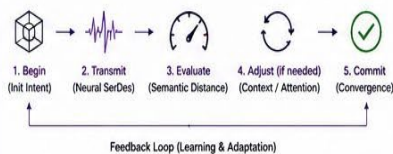
Category	Traditional	Hyper-Aether
What is transmitted?	Bits / Bytes	Semantic State (Meaning)
What is guaranteed?	Exact delivery	Correct understanding
How is success measured?	ACK / Timeout	Semantic distance (D_H)
What causes failure?	Packet loss / Corruption	Misunderstanding / Context mismatch
How is recovery done?	Retransmit packets	Adjust context / Attention / Memory / Resend
What is committed?	Data written	Meaning integrated (Resonant state)
What is rolled back?	Unwritten data	Divergent semantic state
System model	Data-centric system	Cognition-centric system
Consistency model	ACID (Atomicity, Consistency, Isolation, Durability)	SCAR (Semantic, Integrity, Context Coherence, Adaptive Recovery, Resonant Commit)
Ultimate goal	Database correctness	Distributed cognition correctness (Shared understanding)

Structure of a Semantic Transaction

$$T_s = (S_i, C_i, M_i, R_i)$$



Transaction Lifecycle in Hyper-Aether



Commit is decided by semantic alignment.

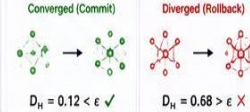
The transaction succeeds only when the receiver understands the same meaning.

Commit Condition

$$D_H(A, B) < \epsilon \rightarrow \text{COMMIT}$$

Where D_H is the semantic distance between sender state (A) and receiver state (B). ϵ is the acceptable semantic divergence threshold.

Semantic Distance Example



From ACID to SCAR

ACID (Traditional)	SCAR (Hyper-Aether)
• Atomicity	• Semantic Integrity
• Consistency	• Context Coherence
• Isolation	• Adaptive Recovery
• Durability	• Resonant Commit

Essence of the Difference

Traditional systems ensure that "the data arrived correctly."

Hyper-Aether ensures that "the meaning was understood correctly."

From data transmission to meaning synchronization.



Hyper-Aether transactions are not about moving data, but about aligning understanding.
The commit moment is the moment of semantic resonance.



各論(VM間通信)

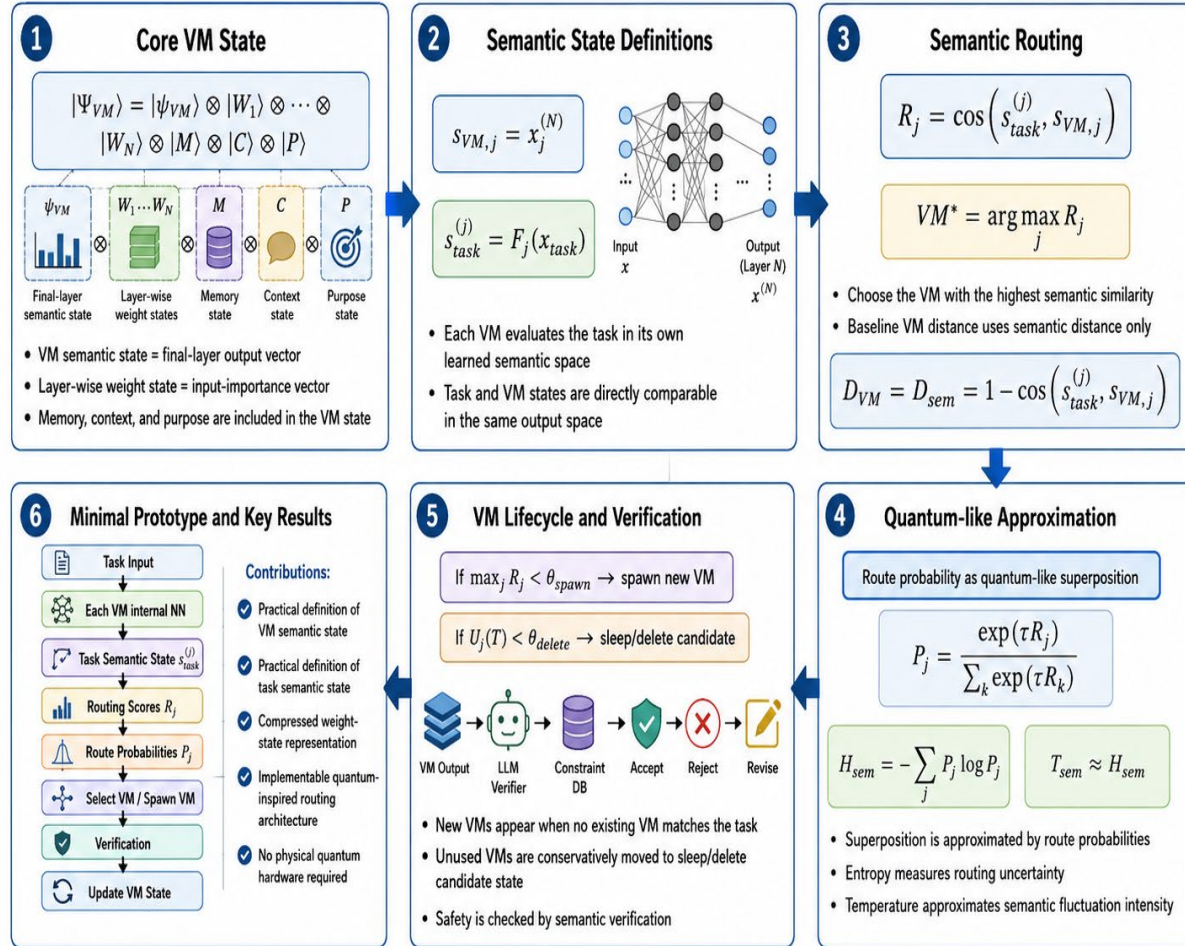
ハイパーエーテル詳細アーキテクチャとコンポーネント説明における意味認知仮想マシン間の神経意味的コミュニケーション

概要

本論文は、Hyper-Aetherフレームワーク内の分散セマンティック認知仮想マシン(VM)向けのセマンティックネイティブ通信アーキテクチャを提案します。従来のシステムでは通信がパケットやシンボル、バイトストリームを転送するのにに対し、Hyper-Aetherは意味認知状態を転送します。このアーキテクチャは以下のものを導入します: • セマンティック・コグニティブ仮想マシン、• ニューラルセマンティックコミュニケーション、• ニューラルシリアルライザー/デシリアルライザー通信、• セマンティックルーティング、• セマンティックメモリー、• セマンティックトランザクション、• 分散認知同期。本論文はさらに各アーキテクチャ要素の詳細な説明を提供し、認知、コミュニケーション、記憶、学習がHyperAether内でどのように相互作用するかを解説します。

Paper Overview: A Practical VM-State Model for Quantum Hyper-Aether

Final-layer semantic states, layer-wise weight states, semantic routing, route entropy, and a minimal prototype



Main Result: Quantum Hyper-Aether can be implemented as a quantum-inspired semantic routing system using ordinary neural networks, cosine similarity, softmax probabilities, entropy-based uncertainty, and verifier-based safety checks.

概要

この図は、量子ハイパーエーテルのための実用的なVM状態モデルに関する論文の概要を示しています。この論文では、抽象的な量子着想に基づく意味論である量子ハイパーエーテルを、最小限の実装可能なアーキテクチャに変換しています。中心となるアイデアは、各VMを、最終層のセマンティック状態、層ごとの重み状態、メモリ、コンテキスト、および目的から構成される結合状態として定義することです。VMのセマンティック状態は、VM内部のニューラルネットワークの最終層出力ベクトルを用いることで具体化されます。タスクのセマンティック状態は、タスク入力を各VMの同じ内部ニューラルネットワークに通すことで定義されます。これにより、各VMは学習済みのセマンティック空間内でタスクを評価できます。セマンティックルーティングは、タスクのセマンティック状態とVMのセマンティック状態間のコサイン類似度によって実行され、ルーティングスコアが最も高いVMが選択されます。基準として、VM間の距離はセマンティック距離のみに単純化されています。この図はまた、物理的な量子ハードウェアを用いずに量子的な振る舞いを近似する方法も示しています。ソフトマックス関数によって生成される経路確率は、量子的な重ね合わせ状態の実用的な近似として使用されます。意味エントロピーは経路確率分布のエントロピーとして定義され、意味温度はこのエントロピーによって近似され、ルーティングの不確実性と意味的変動の強度を表します。VMライフサイクル制御も含まれています。既存のVMがタスクに十分に適合しない場合は新しいVMが生成され、使用されていないVMはスリープ状態または削除候補状態に移行されます。意味検証は、LLM検証器と制約データベースによって実行されます。要約すると、図は、Quantum Hyper-Aetherが、通常のニューラルネットワーク、最終層の意味ベクトル、コサイン類似度、ソフトマックス経路確率、エントロピーに基づく不確実性、VMライフサイクルルール、および検証器に基づく安全性チェックを使用して、量子に着想を得た意味ルーティングシステムとして実装できることを示しています。

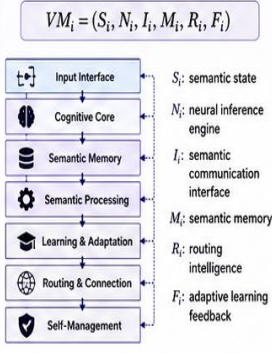
Quantum Hyper-Aether

• Semantic Geometry, Distributed Cognitive Virtual Machines, Semantic Entropy, and Semantic Temperature •

1 Core Idea

- Computation is modeled as semantic graph evolution.
- Virtual machines behave as semantic-cognitive entities.
- Communication becomes meaning transfer and semantic resonance.
- Memory becomes semantic topology.
- Intelligence emerges from distributed semantic interaction.

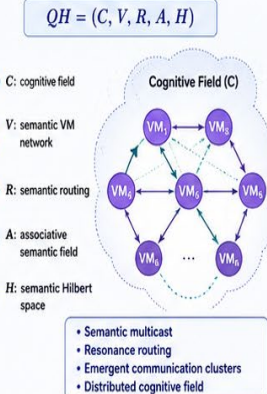
2 Semantic-Cognitive VM



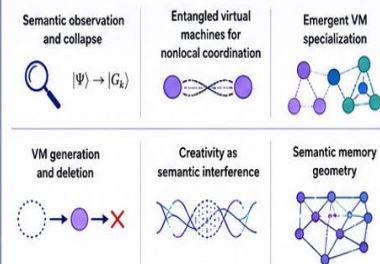
3 Semantic Geometry

- Associative Semantic Fields $A_i = \{a_1^i, a_2^i, \dots, a_m^i\}$
 - Weighted Semantic Similarity $\text{Sim}(i, j) = \frac{|A_i \cap A_j|}{|A_i \cup A_j|}$
 - Quantum Semantic Similarity $\text{Sim}_q(i, j) = |\langle \Phi(S_i) | \Phi(S_j) \rangle|^2$
 - Composite Semantic Distance $D_q(i, j) = 1 - \text{Sim}_q(i, j)$
 $D_G(i, j) = 1 - \text{Sim}_G(i, j)$
 $D_H(i, j) = \alpha D_q + \beta D_G + \gamma D_C$
- Meaning similarity emerges through associative resonance rather than only vector proximity.

4 System Architecture



5 Semantic Dynamics



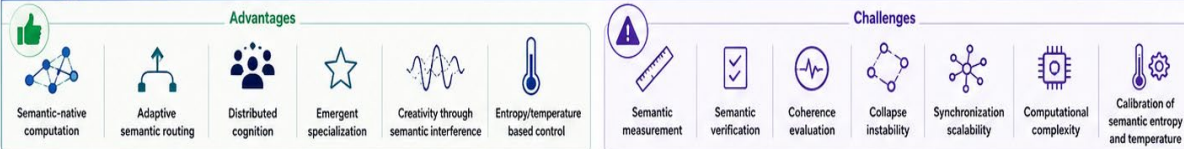
6 Semantic Entropy

- Semantic entropy is the uncertainty, dispersion, and fragmentation of semantic states.
- $H_G = -\sum_i p_i \log p_i$ (state entropy)
 - $H_{VM} = -\sum_k P(C_k) \log P(C_k)$ (cluster entropy)
 - $H_D = \sum_{i,j} p_i p_j D_H(i, j)$ (dispersion entropy)
- Moderate entropy supports creativity and exploration.
• Excessive entropy causes fragmentation and instability.

7 Semantic Temperature

- Semantic temperature is the semantic fluctuation intensity that drives exploration and adaptation.
- $\frac{1}{T_{sem}} = \frac{\partial H_{sem}}{\partial E_{sem}}$ $T_{sem} \approx a \text{Var}(D_H) + b \text{Var}(P_{ij}) + c \text{Var}(S_i) + d R_{collapse} + e R_{birth/delete}$
- Behavior by Semantic Temperature Region
- | Region | Range (T_{sem}) | Behavior |
|-----------------------|---------------------|--|
| Low Temperature | $T_{sem} \ll 1$ | Stable, deterministic, automation |
| Medium Temperature | ≈ 1 | Balanced, adaptive reasoning |
| High Temperature | > 1 | Exploratory, creative, analogical |
| Excessive Temperature | $\gg 1$ | Turbulence, fragmentation, collapse risk |

8 Advantages and Challenges



9 Takeaway

Quantum Hyper-Aether provides a theoretical foundation for AI-native operating systems, distributed AGI infrastructures, semantic computing, semantic thermodynamics, and semantic physics.

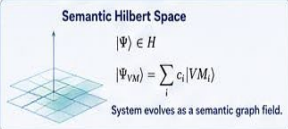
本論文では、セマンティック幾何、分散型認知仮想マシン、セマンティック・エントロピー、セマンティック温度を統合した、意味ネイティブな計算基盤 **Quantum Hyper-Aether** を提案する。本フレームワークでは、計算をセマンティック・グラフの進化として捉え、仮想マシンを意味伝達、セマンティック・ルーティング、分散認知を通じて相互作用するセマンティック認知主体としてモデル化する。さらに、連想的意味距離、量子的意味距離、グラフ意味距離を統合した複合セマンティック幾何を導入し、意味観測、意味収縮、VMの専門化、セマンティック干渉による創造性を記述する。加えて、本論文では、セマンティック・エントロピーを意味状態の不確実性、分散、断片化を表す指標として定義し、セマンティック温度を意味状態の揺らぎの強さを表す指標として導入する。これにより、低温域における自動化、中温域における適応的推論、高温域における創造的探索、過高温域における意味的乱流や収縮不安定性を統一的に説明できる。これらの概念は、システムの安定性、探索性、創造性を制御する理論的基盤を提供する。以上により、Quantum Hyper-Aether は、AIネイティブOS、分散型AGI基盤、セマンティック・コンピューティング、セマンティック熱力学、ならびにセマンティック物理学に向けた理論的基礎を提供する。

Quantum Hyper-Aether: Summary of Concepts and Design Choices

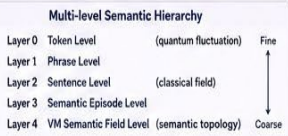
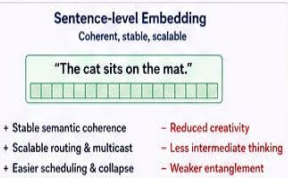
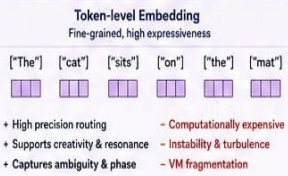
Semantic-Cognitive Virtual Machines • Semantic Geometry • Quantum Semantic Communication

1. FOUNDATIONS

- Computation = Semantic Graph Evolution
- Virtual Machines = Semantic-Cognitive Entities
- Communication = Semantic Resonance
- Memory = Semantic Topology
- Intelligence = Distributed Semantic Interaction



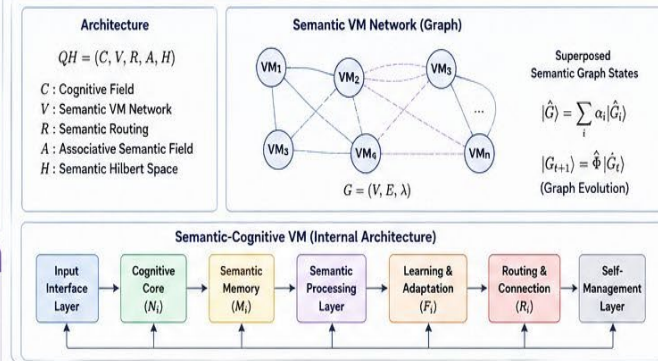
2. SEMANTIC REPRESENTATIONS



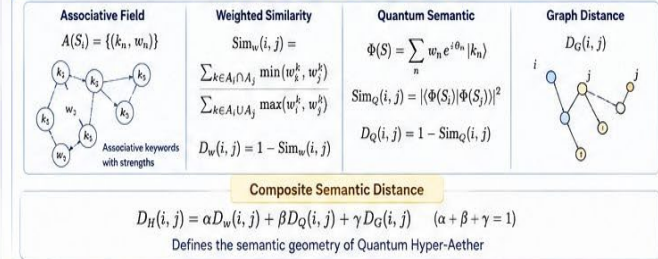
6. KEY OPEN (UNIMPLEMENTED) AREAS

- Semantic Hilbert Space (dimension, basis, phase)
- Embedding Generation (model, granularity, dynamics)
- Associative Field Generator (how to produce keywords)
- Distance Weights (α, β, γ) (learning & adaptation)
- Semantic Graph Runtime (storage, scaling, partitioning)
- VM Internal Models (NN architecture, learning)
- Communication Protocol (format, transport, reliability)
- Semantic Collapse Mechanism (trigger, coherence, rollback)
- Adaptive Learning Loop (objectives, rewards, continual)
- Semantic OS Scheduler (affinity, load balancing)
- Semantic Memory System (storage hierarchy)
- Hardware Implementation (GPU, FPGA, optical, etc.)

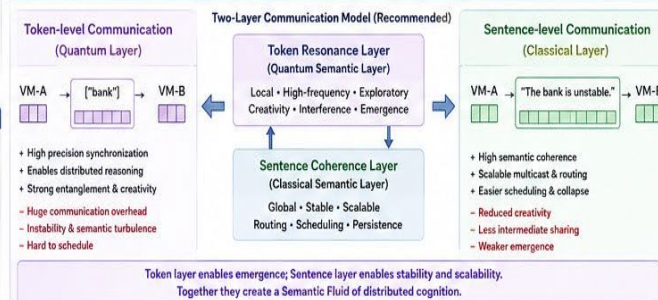
3. QUANTUM HYPER-AETHER ARCHITECTURE



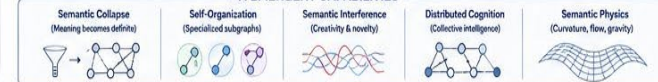
4. SEMANTIC DISTANCE & GEOMETRY



5. VM COMMUNICATION: TOKEN vs SENTENCE



7. EMERGENT CAPABILITIES



8. TOKEN vs SENTENCE: COMPARISON

	Token-level	Sentence-level
Precision	Very High	High (macro)
Creativity	Very High	Moderate
Entanglement	Strong	Moderate
Communication Cost	Very High	Low
Stability	Low	High
Scalability	Low	High
Best For	Exploration, Creativity, Local Resonance	Routing, Scheduling, Stability, OS Layer

Takeaway

Use token-level for emergence and creativity.
Use sentence-level for stability and scalability.
Combine both for the best performance.

9. IMPLEMENTATION ROADMAP (Recommended)



10. ULTIMATE VISION

From Semantic Computing to Semantic Physics

Tokens (Micro) → Sentences (Macro)
→ Meaning Field → Semantic Space-Time

A new computational universe
where meaning itself becomes the fundamental force.



本図は、Quantum Hyper-Aetherを、セマンティック単位で動作する次世代の分散計算基盤として整理したものである。従来のクラウドやOSが、CPU・メモリ・ネットワーク・VMを物理的または論理的な資源として扱うのに対し、本構想では、**意味・文脈・目的・知識状態**を主要な制御対象とする。各VMは単なる実行環境ではなく、認知コア、意味メモリ、学習機構、ルーティング機構を備えた**Semantic-Cognitive VM**として定義される。

中心的な設計要素は、VM群を意味空間上のグラフとして扱う点にある。セマンティック距離は、連想キーワード、重み付き類似度、量子的意味状態、グラフ距離を組み合わせで定義され、VM間通信、ルーティング、スケジューリング、VM生成・削除の判断に利用される。これにより、単なる負荷分散ではなく、**意味的に近い処理を近いVMへ割り当てる**意味ベースの分散制御が可能になる。

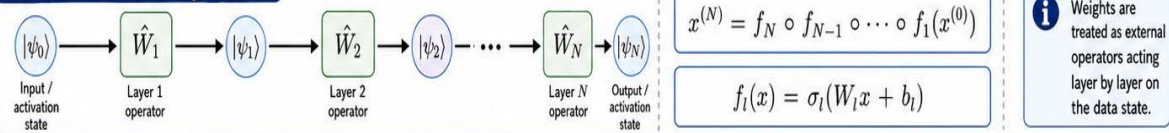
通信方式としては、**token-level communication**と**sentence-level communication**の二層構造が提案されている。tokenレベルは高精度・高創発性・強い意味共鳴を持つが、通信コストが高く不安定になりやすい。一方、sentenceレベルは安定性、スケーラビリティ、OS層での制御に優れる。したがって、実用上は、局所的な創発や探索にはtokenレベルを用い、全体制御、スケジューリング、永続化、ルーティングにはsentenceレベルを用いるハイブリッド構成が最適とされる。

未実装領域としては、セマンティックヒルベルト空間、埋め込み生成器、連想場生成器、距離重み、通信プロトコル、VM内部モデル、セマンティックOSスケジューラ、記憶管理、ハードウェア実装などが挙げられる。実装ロードマップでは、まず埋め込み粒度とモデルを確定し、次に連想場生成器、距離エンジン、セマンティックグラフランタイム、VM通信、スケジューラ、分散実行環境、ハードウェアアクセラレーションへ進む段階的開発が推奨されている。

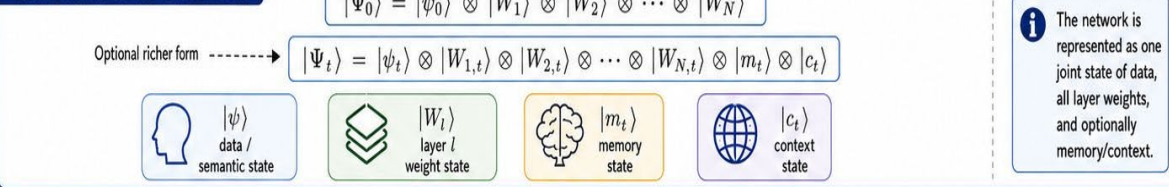
最終的にQuantum Hyper-Aetherは、トークンから文、意味場、セマンティック時空へと拡張される**意味中心の計算宇宙**を目指す。そこでは、計算資源の管理だけでなく、意味そのものの生成、伝播、干渉、崩壊、自己組織化がシステムの基本原理となる。これは、従来の情報処理基盤を超えて、**Semantic ComputingからSemantic Physicsへ進むための理論的・実装的フレームワーク**である。

N-Layer Neural Network with Weights Included in the Quantum-Style State

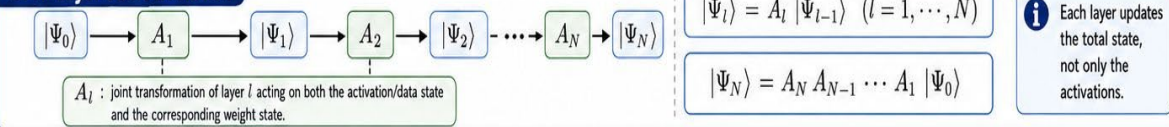
1. Conventional N-Layer View



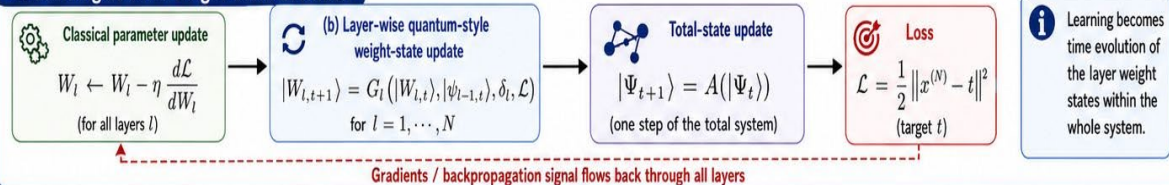
2. Extended Joint-State View



3. N-Layer Joint Evolution



4. Learning as Joint Weight-State Evolution



5. Interpretation and Implications

	Operator-based interpretation	State-inclusive N-layer interpretation
Weights	External operators	Part of the total state
Learning	Operator updates	Weight-state evolution
Inference vs learning	Partly separated	Unified in total-state dynamics
System view	state + operators	joint evolving multilayer system
Physical analogy	effective transformation	open / effective quantum-like system

VM / Semantic Systems View

- semantic state
- layer weight states
- memory
- context
- goal / policy

$|\Psi_{VM}\rangle = |\psi_{sem}\rangle \otimes |W_1\rangle \otimes \dots \otimes |W_N\rangle \otimes |M\rangle \otimes |C\rangle \otimes |P\rangle$



Summary: When weights of all N layers are included in the quantum-style state, an N -layer neural network is interpreted as a joint dynamical system in which activations, layer-wise weights, memory, and context can co-evolve through sequential layer transformations.

この図は、すべての層の重みをシステム状態の一部として含める、NNN層ニューラルネットワークの拡張された量子スタイル解釈を示しています。従来の見方では、各層が外部の重み演算子 W_l を活性化状態に適用し、 $|\psi_0\rangle$ から $|\psi_N\rangle$ への変換列を生み出します。これは標準的なニューラルネットワークの定式化に対応し、推論は層関数の前方適用であり、学習は外部パラメータを更新します。

拡張ビューはネットワークを、データまたはセマンティック状態、すべてのレイヤーごとの重み状態、そしてオプションでメモリおよびコンテキスト状態を含む共同量子スタイルの状態として表現します。この定式化では、全体状態は $A_1, A_2, \dots, A_N$ 、 $\dots$ 、 $A_N A_1, A_2, \dots, A_N$ といった結合層変換の連続を経て進化し、各変換は活性化状態だけでなく関連する重み状態成分も更新します。したがって、NNN層ネットワークは単なる固定オペレーターの連鎖ではなく、多層的な動的システムとして解釈されます。

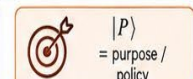
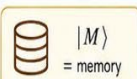
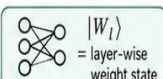
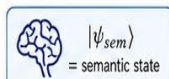
学習もまた再解釈されます。学習を単なる古典的なパラメータ更新として扱うのではなく、 $W_l \leftarrow W_l - \eta \frac{d\mathcal{L}}{dW_l}$ ではなく、図は学習を層ごとの重み状態の時間的変化として記述しています。逆伝播はフィードバック信号となり、全ての状態を伝播し、すべての層にわたる重み状態成分を修正します。

VMやHyper-Aetherのようなセマンティックシステムにおいて、この解釈はVMをセマンティック状態、すべてのレイヤー重み状態、メモリ、コンテキスト、目標またはポリシー状態を含む共同状態としてモデル化できることを示唆しています。まとめると、すべてのNNN層の重みが量子スタイルの状態に含まれると、NNN層ニューラルネットワークは活性化、重み、記憶、文脈が逐次変換を通じて共進化する共同進化システムとなります。

Applying the VM-State Model to Hyper-Aether

1. VM as a Joint Semantic State

$$|\Psi_{VM}\rangle = |\psi_{sem}\rangle \otimes |W_1\rangle \otimes \dots \otimes |W_N\rangle \otimes |M\rangle \otimes |C\rangle \otimes |P\rangle$$



joint evolving VM state



A VM is no longer just code + data + runtime.
It becomes a self-adaptive semantic processing unit.

2. Internal VM Dynamics

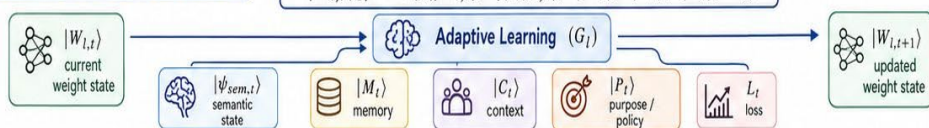
$$|\Psi_{VM}(t+1)\rangle = A_{VM}(|\Psi_{VM}(t)\rangle, |\Phi_{env}(t)\rangle)$$



Inference, memory update, context update, and self-learning are unified in one state-evolution model.

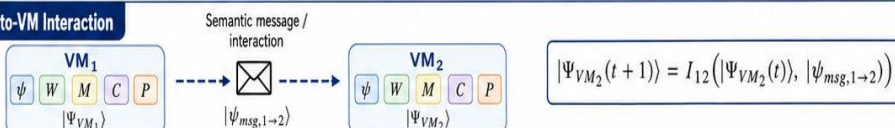
3. Learning as Weight-State Evolution

$$|W_{i,t+1}\rangle = G_i(|W_{i,t}\rangle, |\psi_{sem,t}\rangle, |M_t\rangle, |C_t\rangle, |P_t\rangle, L_t)$$



Learning changes the VM state itself, not only external parameters.

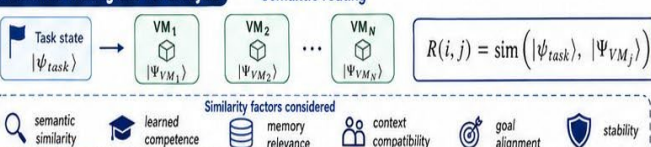
4. VM-to-VM Interaction



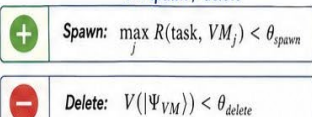
Communication is modeled as state interaction rather than mere packet transfer.

5. Semantic Routing and VM Lifecycle

Semantic routing

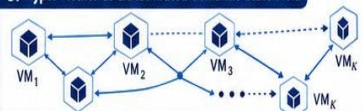


VM spawn / delete



create a new VM when no existing VM fits
remove a VM when its state value becomes too low

6. Hyper-Aether as a Distributed Semantic State Field



$$|\Omega_{HA}\rangle = |\Psi_{VM_1}\rangle \otimes |\Psi_{VM_2}\rangle \otimes \dots \otimes |\Psi_{VM_K}\rangle$$

Hyper-Aether becomes a distributed field of interacting self-adaptive VM states.

System-level functions



Summary:

Applying the VM-state model to Hyper-Aether yields a self-adaptive semantic-state computing fabric. Each VM carries semantic state, layer-wise weights, memory, context, and purpose, while system behavior emerges from the interaction and evolution of many such VM states.

概要

この図は、VM-stateモデルがハイパーエーテルにどのように適用できるかを説明しています。このモデルでは、各VMはもはやコード、データ、ランタイムで構成される従来の実行ユニットとして扱われるのではなく、VMは現在の意味状態、レイヤーごとのニューラルネットワーク重み状態、メモリ、コンテキスト、目的またはポリシーを含む共同セマンティックステートとして表現されます。これにより、各VMは自己適応型セマンティック処理ユニットとして振る舞うことが可能になります。

図はVMの挙動を状態進化として説明できることを示しています。現在のVM状態は環境、ユーザー入力、セマンティックファブリックと相互作用し、推論、適応、メモリ更新、コンテキスト更新、自己学習を通じて次のVM状態へと進化します。学習はまた、意味状態、記憶、文脈、目的、喪失が共に各層の重み状態の更新に影響を与える内部重み状態の進化として再解釈されます。

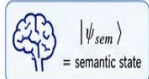
VM-to-VM通信は、単純なパケット転送ではなく、意味状態間の相互作用として説明されます。あるVMからのセマンティックメッセージは受信側VMの状態を変更します。この状態モデルに基づき、セマンティックルーティングは、意味的類似性、学習能力、記憶関連性、文脈互換性、目標整合性、安定性を考慮しながら、タスク状態と候補のVM状態を比較することで、最も適したVMを選択できます。

図はまた、モデルをVMライフサイクル管理にも拡張しています。既存のVMがタスクに十分にマッチしないと新しいVMが生成され、状態値が低すぎると削除されることがあります。システムレベルでは、Hyper-Aetherは多くの相互作用する自己適応型VM状態からなる分散型セマンティック状態フィールドとなります。その結果、通信、ルーティング、学習、VMの作成・削除、意味的調整は、自己適応型セマンティックステートコンピューティングのファブリック内で状態進化プロセスとして統一されることが可能です。

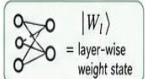
Applying the VM-State Model to Quantum Hyper-Aether

1. QVM as a Quantum-Like Joint Semantic State

$$|\Psi_{QVM}\rangle = |\psi_{sem}\rangle \otimes |W_1\rangle \otimes \dots \otimes |W_N\rangle \otimes |M\rangle \otimes |C\rangle \otimes |P\rangle \otimes |Q\rangle$$



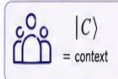
$|\psi_{sem}\rangle$
= semantic state



$|W_i\rangle$
= layer-wise weight state



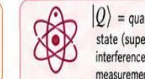
$|M\rangle$
= memory



$|C\rangle$
= context



$|P\rangle$
= purpose/policy



$|Q\rangle$ = quantum-control state (superposition / interference / entanglement / measurement)

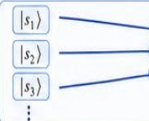
joint evolving quantum-like VM state

i A QVM is a quantum-like semantic processing unit with semantic, weight, memory, context, purpose, and quantum-control states.

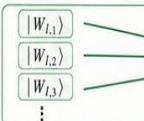
2. Superposition of Meaning and Weight States

$$|\psi_{sem}\rangle = \sum_i \alpha_i |s_i\rangle$$

$$|W_i\rangle = \sum_k \beta_{i,k} |W_{i,k}\rangle$$



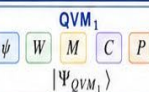
Superposed semantic state
 $|\psi_{sem}\rangle$
 $= \sum_i \alpha_i |s_i\rangle$



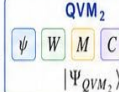
Superposed weight state
 $|W_i\rangle$
 $= \sum_k \beta_{i,k} |W_{i,k}\rangle$

(i) A QVM can hold multiple candidate meanings and multiple transformation modes before selection.

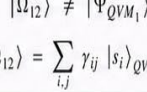
3. Entangled VM-to-VM Semantic Interaction



$|\Psi_{QVM_1}\rangle$



Semantic interaction / entanglement



$|\Psi_{QVM_2}\rangle$

$$|\Omega_{12}\rangle \neq |\Psi_{QVM_1}\rangle \otimes |\Psi_{QVM_2}\rangle$$

$$|\Omega_{12}\rangle = \sum_{i,j} \gamma_{ij} |s_i\rangle_{QVM_1} \otimes |s_j\rangle_{QVM_2}$$

(i) Communication is modeled as entangled semantic interaction rather than simple packet transfer.

4. Quantum Semantic Routing



Task state
 $|\psi_{task}\rangle$

Routing superposition state
 $|\rho_{route}\rangle = \sum_j r_j |route_j\rangle$

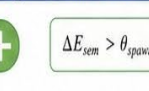
Candidate QVMs
 $QVM_1, QVM_2, \dots, QVM_N$

Semantic similarity & competence
modulate amplitudes
 $r_j \leftarrow r_j \cdot \text{sim}(|\psi_{task}\rangle, |\Psi_{QVM_j}\rangle)$

Measurement / selection
Measure($|\rho_{route}\rangle$) $\rightarrow QVM_k$

(i) Routing is a quantum-like selection process: candidate routes are weighted, interfered, and finally measured.

5. VM Spawn / Delete as Quantum-Like Phase Transition



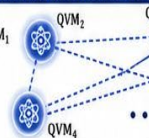
Spawn
 $\Delta E_{sem} > \theta_{spawn}$
 $H(|\rho_{route}\rangle) > \theta_{uncertainty}$
Create a new QVM when the semantic field cannot stably resolve a task with existing QVMs.

Delete
 $A_{QVM} \rightarrow 0$
Remove a QVM when its amplitude, value, or semantic contribution fades out.

Lifecycle management is interpreted as instability, transition, and decay in the semantic field.

(i) Spawn and delete events emerge from phase transitions, instability, and decay in the quantum-like semantic field.

6. Quantum Hyper-Aether as an Entangled Semantic State Field



$|\Omega_{QHA}\rangle = \mathcal{E}(|\Psi_{QVM_1}\rangle, |\Psi_{QVM_2}\rangle, \dots, |\Psi_{QVM_k}\rangle)$

Quantum Hyper-Aether becomes a distributed field of interacting, entangled, self-adaptive QVM states.

System-level functions
communication, routing, learning, spawn/delete, semantic coordination, contextual alignment

(i) The Quantum Hyper-Aether is a distributed field of interacting, entangled, self-adaptive QVM states, which provides a theoretical foundation for the system-level functions.



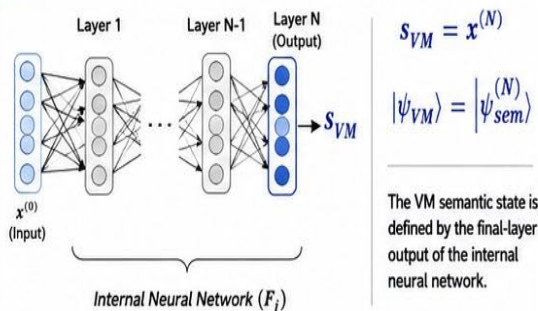
Summary:

Applying the VM-state model to Quantum Hyper-Aether yields an entangled field of self-adaptive quantum-like VM states. Each QVM carries semantic state, layer-wise weights, memory, context, purpose, and quantum-control state, while system behavior emerges from superposition, interference, entanglement, measurement, and co-evolution across the distributed semantic field.

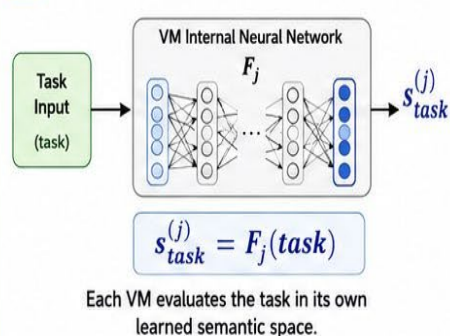
概要 本論文は、量子ハイパーエーテルと呼ばれる拡張された意味論ネイティブ計算フレームワークを提案します。このフレームワークは、意味認知仮想マシン、連想的意味距離、量子に着想を得た意味計算、動的に進化する意味グラフ、神経意味コミュニケーション、分散認知、意味エントロピー、意味温度、そして量子に似た仮想マシン状態モデルを統合しています。本論文で導入された主な拡張は、仮想マシン状態自体に内部ニューラルネットワークの重み付けを含めることです。仮想マシンを単にデータ上でコードを実行するコンテナとして扱うのではなく、提案されたモデルは各VMをセマンティック状態、層ごとのニューラルウェイト状態、メモリ、コンテキスト、目的、ポリシー、量子制御状態を含む共同の意味状態として表現します。したがって、学習はもはや外部パラメータ最適化だけとして解釈されるものではありません。それは自己適応的な意味実体内の内部重み状態の進化の時間となります。VMに埋め込まれたN層ニューラルネットワークの場合、各層の重みは状態成分として表現されます。完全なVM状態は、活性化および重み状態成分の両方を更新するジョイント変換を通じて進化する。量子超エーテルでは、重ね合わせ、干渉、もつれ、測定に似た選択を導入することでこのモデルがさらに拡張されています。VM間の通信は絡み合った意味的相互作用となり、ルーティングは量子的なルート選択となり、VMの生成や削除は意味的相転移として解釈できます。このアーキテクチャは、量子ハイパーエーテルを自己適応型量子様のVM状態からなる絡み合った分散セマンティック状態フィールドとしてモデル化しています。これはAIネイティブOS、セマンティックコンピューティングファブリック、適応ルーティング、自律学習、分散認知、セマンティック熱力学、自己組織化セマンティックインフラストラクチャの理論的基盤を提供します。

Initial Practical Definition of the Quantum Hyper-Aether VM-State Model

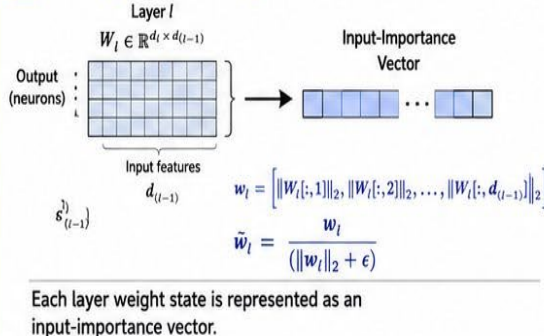
1. VM Semantic State



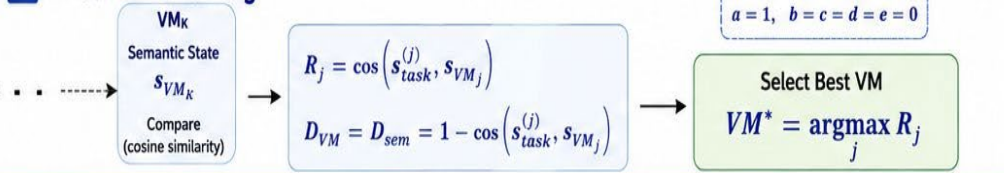
2. Task Semantic State



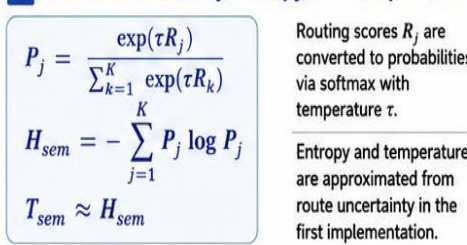
3. Layer-wise Weight State



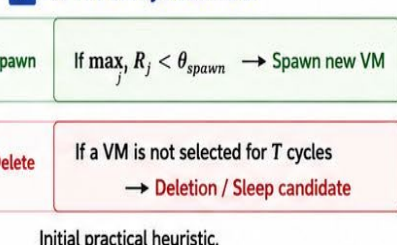
4. Semantic Routing



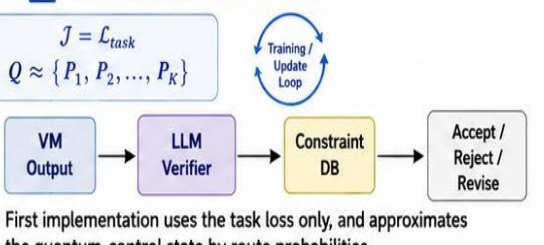
5. Route Probability, Entropy, and Temperature



6. VM Lifecycle Control



7. Learning and Verification



Recommended First Prototype

- Minimal working model
- Same internal NN for VM and task encoding
- Final-layer semantic state
- Weight state as input-importance vector
- Routing by cosine similarity
- Entropy-based uncertainty
- LLM verifier + Constraint DB

概要

この図は、量子ハイパーエーテルVM-stateモデルの初期の実用的定義をまとめたものです。このモデルは、VMのセマンティック状態をVMの内部ニューラルネットワークの最終層出力として定義し、セマンティック状態をベクトルとして直接測定可能にしています。タスクの意味状態は、各VMの同じ内部ニューラルネットワークを通じてタスク入力を通すことで得られ、各VMは自分の学習した意味空間内でタスクを評価します。

図はまた、層ごとの重み状態を入力重要度ベクトルとして定義しています。重み行列全体を表すのではなく、各層の重み状態はベクトルに圧縮され、その成分は各入力次元がその層にどれだけ強く寄与するかを示します。これにより、重み状態の進化を実用的かつ解釈可能な表現として表現できます。

セマンティックルーティングは、タスクのセマンティック状態と各VMのセマンティックステートとのコサイン類似性によって定義されます。初期の基準として、VM距離は意味距離のみに簡略化され、 $a=1$ および $b=c=d=e=0$ という設定を用います。最適なVMは最も高いルーティングスコアで選ばれ、ルーティング確率はソフトマックス分布を用いて計算されます。このルーティング不確実性から意味的エントロピーと意味温度が近似されます。

図はまた、単純なライフサイクルルールを紹介しています。既存のVMがタスクと十分な意味的類似性を持っていない場合に新しいVMが生成され、未使用VMは一定期間後に削除またはスリープ候補となります。最初の実装では、自律学習はタスク損失のみに簡素化され、量子制御状態はルート確率で近似され、意味検証はLLM検証器と制約データベースによって行われます。

まとめると、図はQuantum Hyper-Aether VM状態モデルの最小限の動作プロトタイプを示し、実践的な実装を強調しています:最終層の意味状態、共有されたVM/タスクエンコーディング、入力重要度の重み状態、余弦ベースのルーティング、エントロピーに基づく不確実性、単純なVMライフサイクル制御、検証者ベースの安全性チェックです。

Simulation Evaluation Results: Sensitivity to θ_{spawn}

Quantum Hyper-Aether Prototype

① Evaluation Setup

- 500 tasks processed
- Initial VM count: 5
- Semantic routing by cosine similarity
- Spawn rule: spawn if $\max R_j < \theta_{spawn}$
- Forced assignment: existing VM selected even though $\max \text{score} < 0.60$
- Compared θ_{spawn} values: 0.30, 0.40, 0.50, 0.60, 0.70

② Summary Table

θ_{spawn}	Final VM Count	Spawn Count	Forced Assignments	Normal Assignments	Avg Max Score	Avg Semantic Entropy
0.30	33	28	266	206	0.489	1.901
0.40	72	67	245	188	0.505	2.777
0.50	137	132	181	187	0.521	3.438
0.60	285	280	0	220	0.501	4.021
0.70	455	450	0	50	0.472	4.385

★ Best balance in this experiment

④ Key Findings



Low θ_{spawn} (0.30–0.40): fewer new VMs, but many forced assignments.



Mid θ_{spawn} (0.50): balanced tradeoff between VM growth and semantic fit.



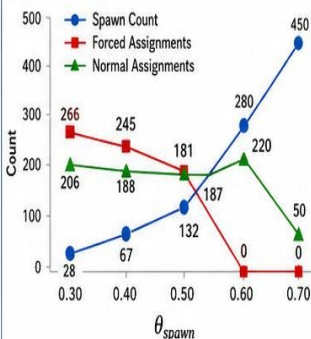
High θ_{spawn} (0.60–0.70): almost no forced assignments, but excessive VM spawning.



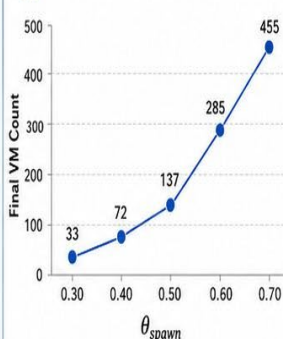
Recommended initial setting:
 $\theta_{spawn} \approx 0.50$.

③ Trend Highlights

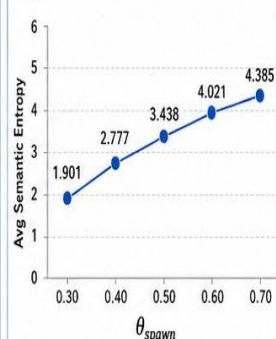
A Assignment Tradeoff



B Final VM Count



C Average Semantic Entropy



Conclusion: Lower θ_{spawn} tends to force tasks into existing VMs, while higher θ_{spawn} over-fragments the semantic space. In this simulation, $\theta_{spawn} = 0.50$ provided the most practical balance.

概要

この図は、VMのスポン閾値

θ_{spawn} に関する量子ハイパーエーテルプロトタイプの実験評価をまとめたものです。この実験は、5台のVMからなる初期プールで500のセマンティックタスクを処理し、異なる閾値がVMの成長、強制割り当て、通常割り当て、ルーティング品質、セマンティックエントロピーにどのように影響するかを評価します。結果は明確なトレードオフを示しています。低い θ_{spawn} 値はVM作成を減らすが、多くの意味的に弱いタスクを既存のVMに強制的に割り当てる一方で、高い

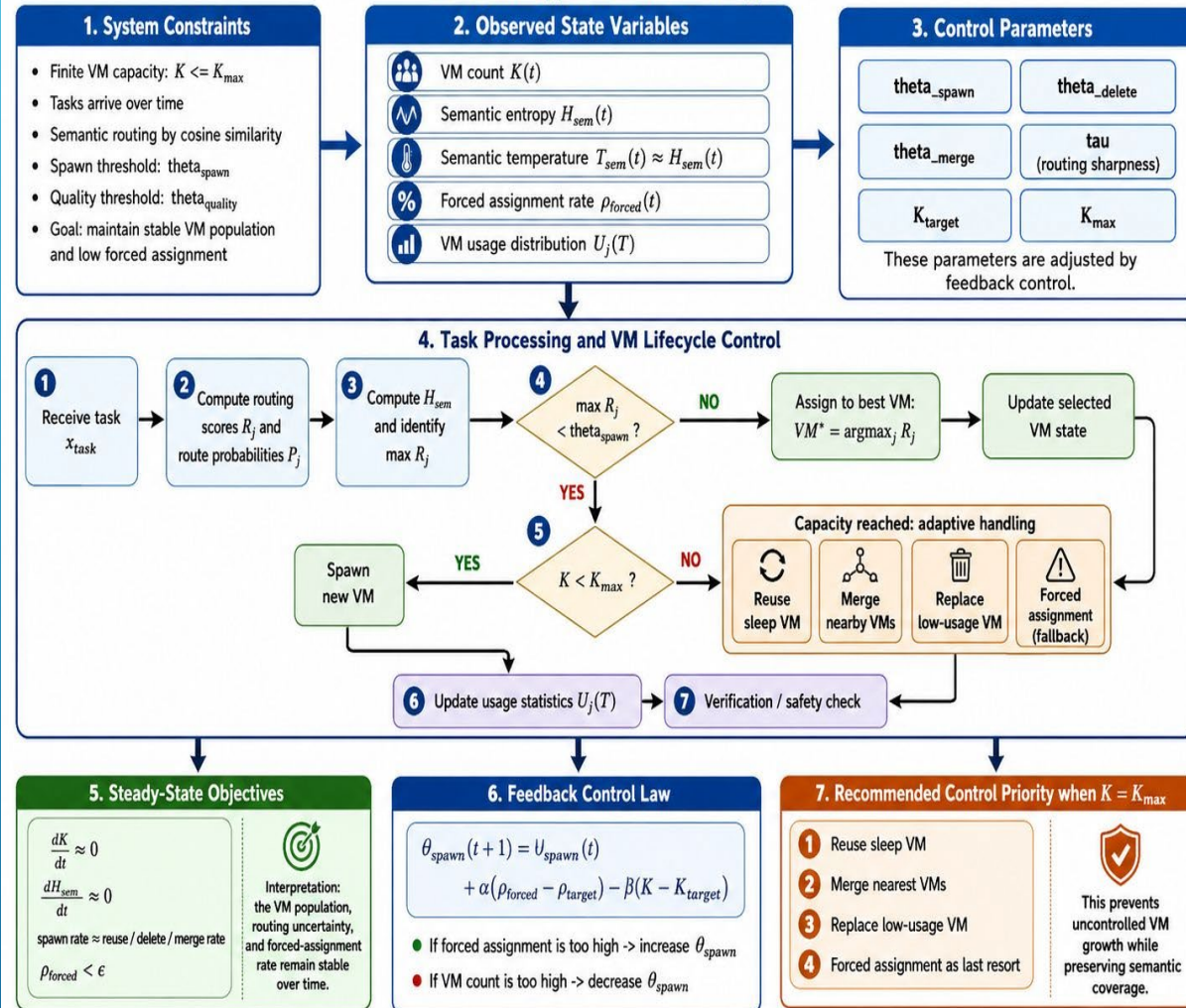
θ_{spawn} 値は強制割り当てを排除しますが、過剰なVM生成と意味の過剰断片化を引き起こします。テストされた数値の中で、

$\theta_{spawn} = 0.50$ が最も実用的なバランスを提供し、適度なVM成長を維持しつつ強制割り当てを減らします。したがって、図は

$\theta_{spawn} \approx 0.50$ をQuantum Hyper-Aetherにおける適応的セマンティックルーティングの初期プロトタイプ設定として使用することをサポートします。

System Control for Finite VMs and Steady-State Maintenance

Quantum Hyper-Aether Prototype



Finite-VM Quantum Hyper-Aether becomes a feedback-controlled semantic space management system.

概要

この図は、Quantum Hyper-Aether有限VMセマンティックルーティングプロトタイプにおけるスケーラビリティに必要な条件をまとめたものです。中心的な主張は、スケーラビリティは単にVMの数を増やすだけでなく、セマンティックルーティングの精度、バランスの取れたVM利用率、限定されたセマンティックエントロピー、低い強制割り当て率、管理可能な検証オーバーヘッドを維持しつつ、効果的なスループットを向上させることで達成されるということです。

この図は、定常状態スループットをアクティブなVM数、1台あたりの平均トランザクションレート、そして全体的な効率率にほぼ比例するものと定義しています。次に、スケーラブルな動作に必要な条件を特定します。タスクは意味的に適切なVMに割り当てられ、ワークロードは過度な集中なしに分散されなければならず、入カトランザクション率はアクティブなVM集団の総サービス容量を下回っている必要があります。

図はまた、良いスケーリングと悪いスケーリングを対比しています。良いスケーリングとは、タスクとVMのマッチングが意味があり、VM間で実行が並列化され、ルーティングオーバーヘッドが低く、再試行やリビジョン率が低く、有限の容量がVMの再利用、マージ、または置き換えによって処理される場合に起こります。多くのVMが冗長化したり、いくつかのVMがボトルネックになったり、強制割り当てが増加したり、セマンティックエントロピーが高くなったり、LLM検証がボトルネックになったりすると、スケーラビリティは崩壊します。

結論として、Quantum Hyper-Aetherは制御された条件下でスケーラブルである。すなわち、セマンティックルーティングの正確性、VM利用のバランス、強制割り当ての低さ、検証オーバーヘッドの制御、フィードバック制御が有限なVM人口をほぼ一定状態に保つ必要がある。

When Required VM Count Becomes Logarithmic in Transaction Volume

Quantum Hyper-Aether Semantic Routing Prototype



1. Core Result

$$N_{VM}(M) \approx N_0 + a \log M$$

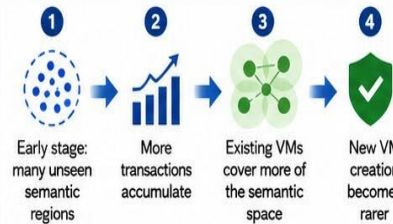
Required VM count grows logarithmically when semantic coverage improves as transactions accumulate.

$$\Pr(\max_j R_j < \theta_{spawn}) \propto \frac{1}{M}$$

The probability of needing a new VM decreases roughly inversely with transaction count.



2. Intuition



At first, new VMs are created frequently; later, most tasks are absorbed by existing VMs.



3. Why Logarithmic Growth Appears

$$\begin{aligned} \Delta N_{VM}(M) &\propto \frac{1}{M} \\ N_{spawn}(M) &= \sum_{m=1}^M \frac{a}{m} \\ &\approx a \log M \end{aligned}$$

Therefore, required VM count grows slowly even as transaction volume increases.



4. Conditions That Enable Log Scaling

- ✓ Stable semantic task distribution
- ✓ Semantic cluster popularity is skewed (frequent patterns repeat)
- ✓ VM states learn cluster centers over time
- ✓ θ_{spawn} is tuned appropriately
- ✓ Similar tasks are absorbed by existing VMs
- ✓ Merge / reuse / replace mechanisms are available
- ✓ Forced assignment remains low



These conditions allow semantic coverage to improve faster than semantic novelty.



5. When Log Scaling Fails

- ! Every transaction introduces nearly new semantics
- ! θ_{spawn} is too high
- ! VM learning is weak or absent
- ! Semantic distribution keeps drifting
- ! No merge / reuse / replacement control
- ! Forced assignment hides quality loss

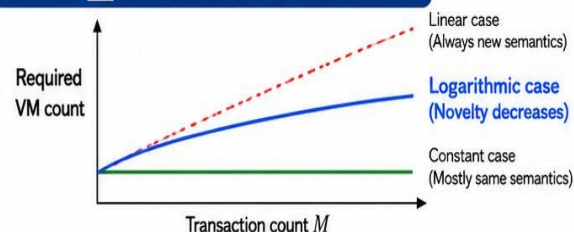


In these cases, VM demand can approach linear growth.

$$N_{VM} \propto M$$



6. Qualitative Growth Patterns



7. Practical Interpretation for Quantum Hyper-Aether

- ! More transactions do not necessarily require proportionally more VMs.
- ! If semantic routing improves coverage, VM growth slows down.
- ! The system becomes more scalable when common semantic regions are reused.
- ! Good control of θ_{spawn} and VM lifecycle is essential.



Target regime: logarithmic VM growth with low forced assignment and stable semantic coverage.



Conclusion: Required VM count becomes approximately *logarithmic* in transaction volume when the probability of discovering a truly new semantic region decreases as $1/M$ and the VM pool progressively learns and covers the task space.

概要

この図は、量子ハイパーエーテルに必要なVM数がトランザクション量とともに対数的に増加する条件を説明しています。重要な結果は、真に新しい意味領域を発見する確率が $1/M$ に減少した場合、VMカウントは $N_{vm}(M) \approx N_0 + a \log M$ として近似できるということです。この体制では、トランザクションがセマンティック・ノウティよりも速く蓄積されるため、既存のVMがタスク空間の拡大する部分をカバーできるようになります。

図は、対数的なVM成長がセマンティックタスクの分布が安定し、頻繁なセマンティックパターンが繰り返され、VMの状態が時間とともにクラスターセンターを学習し、 $\theta(\text{spawn})$ が既存のVMに吸収されるように調整されていることを示しています。マージ、再利用、置換の仕組みは、過剰なVMの増加をさらに抑制しつつ、強制割り当てを低く抑えます。

図はまた、対数成長と故障のケースを比較しています。すべてのトランザクションがほぼ新しい意味論を導入し、 $\theta(\text{spawn})$ が高すぎる場合、VM学習が弱い場合、または意味分布がドリフトし続ける場合、必要なVM数は線形成長に近づくことがあり、 $N_{vm} \propto M$ に近づくことができます。したがって、対数スケーリングは自動的に行われるものではありません。安定したセマンティックカバレッジと効果的なVMライフサイクル制御に依存します。

結論として、Quantum Hyper-Aetherはシステムが共通の意味領域を段階的に学習・再利用することでスケーラブルなVM成長を実現でき、トランザクション量の増加に伴い新たなVM作成がますます稀になることが示されています。

Finite-VM Simulation Results: Steady State, Throughput, and Logarithmic VM Growth

Quantum Hyper-Aether Prototype

1

Evaluation Setup

- 3,000 transactions processed
- Finite VM capacity K_{max} tested at 50, 100, 150
- Initial VM count: 5
- Semantic routing with cosine similarity
- Feedback control on θ_{spawn}
- Capacity handling: reuse sleep VM, merge nearby VMs, replace low-usage VM, forced assignment as last resort
- Metrics: steady state, throughput, forced assignment, log VM growth

2

Summary Results

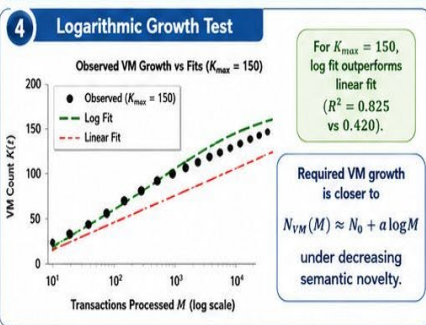
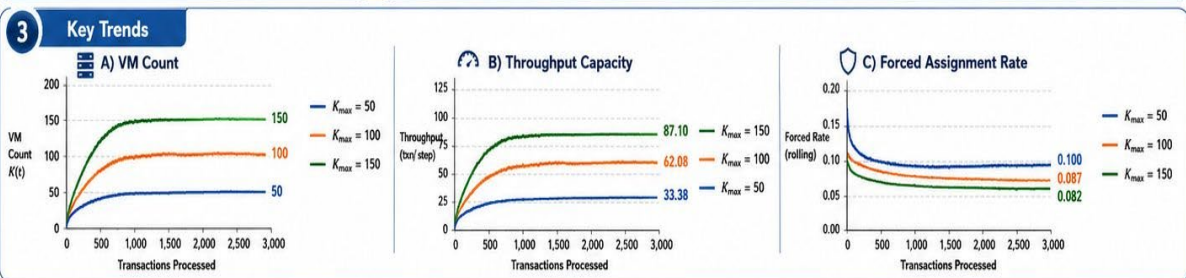
K_{max}	Final VM Count	Avg VM Count (steady)	Forced Rate (steady)	Avg Throughput (steady)	Log Fit R^2	Linear Fit R^2
50	50	50.0	0.100	33.38	0.421	0.086
100	100	100.0	0.087	62.08	0.798	0.361
150	150	150.0	0.082	87.10	0.825	0.420

✓

The $K_{max} = 150$ case shows the strongest evidence for log-scaling.

i

All cases reached $K(t) = K_{max}$ and then stabilized at steady state.



5

Main Findings

✓

Steady State: After saturation, $K(t)$ remains stable at K_{max} .

📈

Scalability: Increasing K_{max} increases steady-state throughput.

🛡️

Quality: Forced assignment remains low and decreases slightly with more VM capacity.

🔗

Structure: The observed growth pattern is more consistent with logarithmic than linear VM growth.

6

Interpretation

⚙️

Finite VM systems can maintain steady operation under feedback control.

📈

Throughput improves with available VM capacity, though not perfectly linearly.

🔄

Semantic reuse reduces the need for continuously spawning new VMs.

🎯

The target regime is stable semantic coverage with low forced assignment.

🏆

Conclusion: In this simulation, finite-VM Quantum Hyper-Aether reaches steady state, scales throughput with K_{max} , and shows VM growth that is closer to logarithmic than linear under stable semantic task distributions.

★

Final θ_{spawn} converged near 0.25–0.26 across the tested capacities.

概要

この図は、Quantum Hyper-Aetherプロトタイプの有界VMシミュレーション結果をまとめたもので、定常状態の挙動、スループットのスケラビリティ、強制割り当て、対数的なVM成長を評価しています。シミュレーションは、5台のVMから始まり、有界のVM容量制約で $K_{max}=50, 100, 150$ の3,000トランザクションを処理し、セマンティックルーティング、 θ_{spawn} のフィードバック制御、VMの再利用、マージ、置換、強制割り当てなどのライフサイクルメカニズムを最終手段として用います。

結果は、すべてのテスト済み構成が最終的に容量上限 $K(t)=K_{max}$ に達し、その後安定することを示しており、有界VM操作がフィードバック制御下で定常状態に入ること示しています。 K_{max} が増加すると推定定常状態スループットも増加し、 $K_{max}=50$ の33.38から $K_{max}=100$ の62.08、 $K_{max}=150$ の87.10に増加します。これにより、利用可能なVM容量の増加が処理能力を向上させることは確認されていますが、完全に線形的ではありません。

シミュレーションはまた、強制割り当て率が比較的低く、VM容量の増加に伴わずかに減少していることを示しています。

$K_{max}=150$ の場合、定常強制割り当て率は0.082であり、VM容量が大きいことで意味カバレッジが向上し、意味的不一致が減少することが示唆されています。さらに、VMの成長パターンは線形フィットよりも対数フィットで説明しやすい。特に $K_{max}=150$ の場合、対数フィットは $R^2=0.825$ を達成し、線形フィットは $R^2=0.420$ である。

結論として、安定した意味的タスク分布の下では、有界VM量子ハイパーエーテルは定常状態を維持し、VM容量に伴うスループットを拡大し、強制割り当てを抑制し、VMの成長を線形よりも対数に近い速度で示すことができます。これは、セマンティック再利用やフィードバック制御型VMライフサイクル管理がスケラブルな運用の重要なメカニズムであるという解釈を支持します。

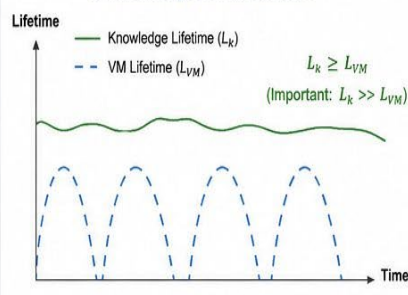
Knowledge Lifetime and Preservation under VM Deletion in Quantum Hyper-Aether

Principle: Knowledge Lifetime $\geq$ VM Lifetime (Important Knowledge: $\gg$ VM Lifetime)

1. Types of Knowledge in a VM

Type	Description	Lifetime (Typical)	What Happens on VM Deletion?
Ephemeral (Short-term)	Recent context, short-term state	Seconds – Minutes	Can be lost (No need to keep)
Learned (Mid-term)	Semantic patterns, weights, usage statistics	Minutes – Hours/Days	Extract & save before deletion
Verified (Long-term)	Verified knowledge (rules, constraints, facts)	Days – Long-term	Persist in external knowledge store

2. Knowledge vs. VM Lifetime



3. Knowledge Loss Model on VM Deletion

$$K_{lost} = K_{VM} - K_{externalized}$$

K_{VM} : Total knowledge in the VM

$K_{externalized}$: Knowledge extracted and saved externally

K_{lost} : Knowledge lost when VM is deleted

➔ Goal: Maximize $K_{externalized}$ before deletion to minimize knowledge loss.

4. Knowledge Lifetime Determined by Importance

Importance Score of Knowledge I_k

$$I_k = a \cdot usage + b \cdot verification + c \cdot novelty + d \cdot stability$$

High Importance $I_k \uparrow$

- ✓ Frequently used
- ✓ Verified
- ✓ Novel
- ✓ Stable & reliable

Knowledge Lifetime
 $L_k = L_0 + \gamma I_k$

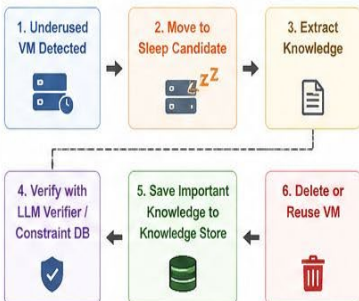
Low Importance $I_k \downarrow$

- ✓ Rarely used
- ✓ Unverified
- ✓ Possibly incorrect
- ✓ Outdated

Shorter Lifetime $L_k \downarrow$

Longer Lifetime $L_k \uparrow$

5. Safe VM Deletion Procedure (Knowledge Preservation First)



6. Hierarchical Knowledge Lifetime Model

Layer	Typical Lifetime	Content / Role
L1 Short-term Context	Seconds – Minutes	Task session context, temporary states
L2 VM-local Knowledge	Minutes – Hours/Days	Local semantic patterns, weights, usage within VM
L3 Shared Semantic Memory	Days – Months	Reusable semantic clusters across VMs
L4 Verified Knowledge Base	Long-term / Semi-permanent	Verified rules, constraints, stable knowledge

7. TTL (Time-To-Live) for Knowledge

$$TTL_k = f(usage, confidence, verification, age, novelty)$$

Usage $\uparrow$

Extend TTL

Verified $\uparrow$

Extend TTL

Novelty $\uparrow$

Extend TTL

Age $\uparrow$

Shorten TTL

Contradiction $\uparrow$

Shorten TTL

TTL expiration $\rightarrow$ Knowledge becomes archive or deletion candidate.

8. Safe Deletion Condition

✓ Usage is low:

$$U_j(T) < \theta_{delete}$$

✓ All critical knowledge is externalized:

$$K_{critical,j} \subseteq K_{external}$$

➔ Only then $\rightarrow$ VM can be deleted or reused.

9. Effects of Proper Knowledge Lifetime Management

✓ Prevents important knowledge loss

✓ Maintains system performance and accuracy

✓ Enables VM reuse and efficient resource control

✓ Keeps knowledge base compact and high-quality

Key Takeaway

VMs are short-lived execution units; knowledge is a long-lived asset.

Externalize, verify, and preserve important knowledge before VM deletion.

Knowledge Lifetime $\geq$ VM Lifetime

💡 By decoupling knowledge lifetime from VM lifetime and managing it with importance, verification, and TTL, Quantum Hyper-Aether can scale with finite VMs while preserving valuable knowledge.

概要

この図は、VMが削除、再利用、または置き換えられた際の Quantum Hyper-Aetherの知識寿命管理モデルをまとめたものです。中心的な原則は、特に重要または検証済みの知識において、知識の寿命はVMの寿命よりも長くすべきだということです。このモデルでは、VMは短命な実行ユニットかもしれませんが、貴重な知識を抽出し、検証し、外部化し、VMを削除する前に保存しなければなりません。この図は、VM知識を短期コンテキスト、VM局所学習知識、共有意味記憶、検証済み長期知識に分類しています。短期的なコンテキストはVMが削除されると安全に消えるかもしれませんが、学習した意味パターン、重み状態情報、使用統計、検証済みルール、制約は外部のナレッジストアに保存されるべきです。知識損失は、総VM知識と外部に保存された知識の差としてモデル化されており、削除の安全性はVM除去前に知識の外部化を最大化することに依存しています。

この図はまた、利用、検証、新規性、安定性に基づく知識重要度スコアを定義しています。重要度の高い知識はより長い寿命を持ち、使用頻度が低い、検証されていない、時代遅れ、または矛盾する知識はより短い寿命を持ちます。TTLベースの仕組みが提案されており、頻繁に使われ検証された知識を拡張し、旧式または信頼性の低い知識はアーカイブや削除候補となります。

結論として、VM削除は単なるリソースグリーンアップとして扱うべきではないということです。代わりに、安全な削除手順に従うべきです。すなわち、使われていないVMを検出し、スリープ候補に移動させ、知識を抽出し、LLM検証器と制約データベースで検証し、重要な知識を外部に保存し、その後にVMを削除または再利用するという手順を踏むべきです。これにより知識寿命とVMの寿命が切り離され、有限VM量子ハイパーエーテルが貴重な知識を保持しつつスケーリングが可能になります。

Knowledge Lifetime Management in Finite-VM Quantum Hyper-Aether

Simulation Results: Knowledge Loss, Reuse, and Performance under VM Reuse

1 Evaluation Setup

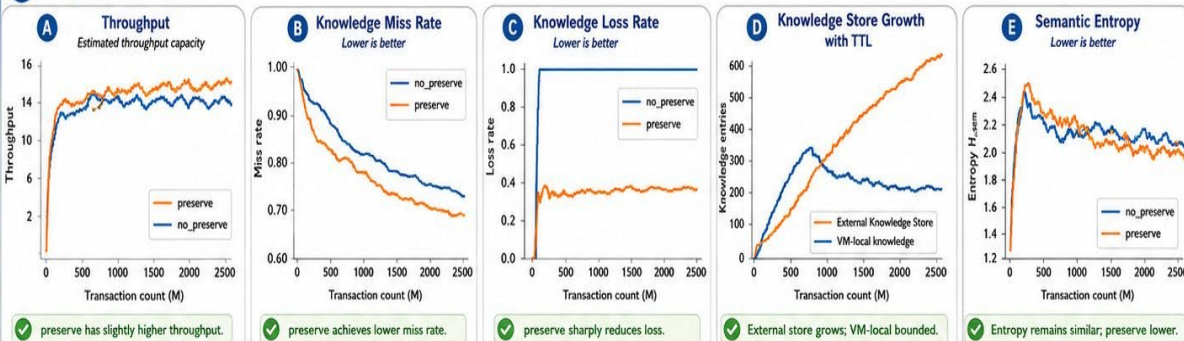
- 2,500 transactions processed
- Finite VM capacity: $K_{\max} = 25$
- Initial VM count: 5
- High novelty task stream to force VM reuse
- Comparison of two modes: "no_preserve" vs "preserve"
- preserve mode uses: external Knowledge Store, importance-based extraction, TTL-based retention

2 Summary Results Table

	Final K	VM reuse events	Forced rate (steady)	Avg entropy (steady)	Avg throughput (steady)	Knowledge loss count	Externalized knowledge count	Knowledge reuse count	Knowledge miss rate	Knowledge loss rate
no_preserve	25	948	0.105	2.059	14.63	1154	0	0	0.729	1.000
preserve	25	1079	0.109	2.011	14.73	465	798	149	0.694	0.368

★ preserve mode (green) is the preferred method.

3 Key Trend Panels



4 Main Findings

- Knowledge preservation sharply reduces loss during VM reuse.
- Externalization + TTL enables reusable long-lived knowledge beyond VM lifetime.
- Knowledge miss rate decreases with preservation.
- Throughput improves slightly with preservation.
- System behavior remains stable under finite VM capacity.

5 Interpretation

- In **no_preserve** mode, VM-local knowledge disappears when a VM is reused, causing total knowledge loss.
- In **preserve** mode, important knowledge is extracted, verified, stored externally, and reused later.
- The result supports the design principle: "Knowledge lifetime $\geq$ VM lifetime".

Legend

- VM: Virtual Machine
- Reuse: VM reuse event
- Knowledge Store: External knowledge repository
- TTL: Time-To-Live (retention window)
- Preserve: Preserve mode (knowledge lifetime management)

概要

この図は有限VM量子ハイパーエーテルシステムにおける知識寿命管理のシミュレーション結果をまとめたものです。この実験では、積極的なVM再利用の2つのモードを比較します。知識保存なしのベースラインモードと、重要度ベースの抽出とTTLベースの保持を用いて重要なVMローカル知識をナレッジストアに外部化する保存モードです。結果は、保存がなければVMのローカル知識が再利用されるたびに失われ、知識損失率は1.000に達していることを示しています。知識保存が有効になると、システムは798の知識エントリを外部化し、そのうち149件を再利用して知識損失率を0.368に低減します。これにより、知識寿命とVM寿命を分離することで、有限のVM容量下で知識損失を大幅に減少させることが確認されます。シミュレーションでは二次的な性能向上も示されています。知識ミス率は0.729から0.694に減少し、意味エントロピーはやや低く、推定定常状態のスループットは14.63から14.73に改善されました。スループットの増加は控えめですが、知識損失やミス率の減少はVM再利用時の意味的連続性の向上を示しています。結論として、知識ライフタイム管理により、有限VM量子ハイパーエーテルは、VMが削除、再利用、置き換えられた場合でも貴重な知識を保持できる。TTL制御で重要な知識を外部化・検証・保持することで、システムは安定した動作を維持しつつ知識損失を減らし、性能をわずかに向上させることができます。

Finite-VM Simulation with Knowledge Lifetime Management

Comparative results for TTL and knowledge verification



EXPERIMENT SETUP

- Synthetic discrete-event simulation
- 4 conditions: A No TTL / No verification; B TTL / No verification; C No TTL / Verification; D TTL / Verification
- 20 maximum VMs, 60 topics
- Knowledge capacity: 12 entries per VM
- 6,000 simulation steps, 12 random seeds
- Goal: evaluate performance and incorrect-knowledge retention under finite VM resources



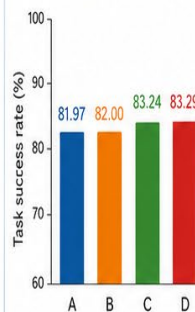
KEY FINDINGS

- ✓ Knowledge verification produced the largest gain in task success.
- ✓ TTL alone had little effect on success rate, but reduced stale incorrect knowledge.
- ✓ The combined condition D achieved the best overall balance.
- ✓ Compared with A, condition D improved success rate by about **1.3 percentage points**.
- ✓ Compared with A, condition D reduced retained incorrect knowledge by about **89%**.
- ✓ Forced assignment rate also decreased under verification.

COMPARISON OF CONDITIONS

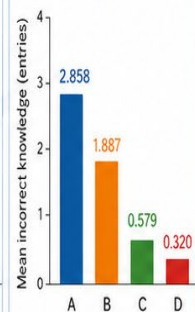
Condition	Task Success Rate (%)	Mean Latency (steps)	Mean Incorrect Knowledge (entries)	Forced Assignment Rate (%)
A No TTL / No verification	81.97	3.247	2.858	5.83
B TTL / No verification	82.00	3.246	1.887	5.74
C No TTL / Verification	83.24	3.295	0.579	4.43
D TTL / Verification	83.29	3.296	0.320	4.34

1 Task Success Rate by Condition



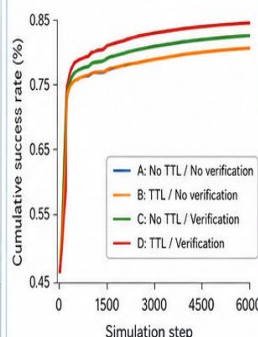
Verification increases success; TTL alone has minimal impact.

2 Retained Incorrect Knowledge (Lower is Better)



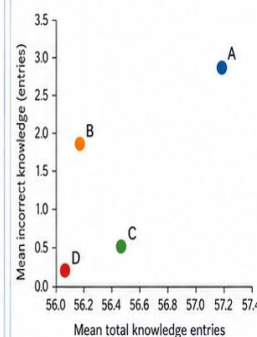
TTL and verification sharply reduce stale incorrect knowledge.

3 Cumulative Success Rate over Time



All conditions converge near ~0.82~0.84; verification (especially D) stays highest.

4 Knowledge Retention vs. Incorrect Knowledge



Condition D achieves the best balance: high knowledge retention with minimal incorrect knowledge.



Separating VM lifetime from knowledge lifetime, and combining TTL with knowledge verification, can improve reliability in finite-VM environments while sharply reducing incorrect-knowledge accumulation.

アブストラクト

本研究では、有限数の仮想マシン（VM）で構成される環境において、知識ライフタイム管理の有効性を合成離散事象シミュレーションにより評価した。比較対象として、TTL（Time-to-Live）なし・知識検証なし、TTLあり・知識検証なし、TTLなし・知識検証あり、TTLあり・知識検証あり、の4条件を設定した。シミュレーションでは、最大VM数を20、意味トピック数を60、各VMの知識保持容量を12件とし、6,000ステップ、12種類の乱数系列で評価を行った。

結果として、知識検証はタスク成功率の向上に最も大きく寄与し、成功率は81.97%から83.24%へ向上した。一方、TTL単独では成功率への影響は小さかったが、保持される誤知識を削減する効果が確認された。TTLと知識検証を併用した条件では、成功率83.29%、平均誤知識数0.320件、強制割当率4.34%となり、総合的に最良の結果を示した。基準条件と比較すると、成功率は約1.3ポイント向上し、保持誤知識は約89%削減された。以上より、VMの寿命と知識の寿命を分離し、TTLによる知識失効と知識検証を組み合わせることで、有限な計算資源下でも信頼性を向上させ、誤知識の蓄積を抑制できる可能性が示された。ただし、本結果は合成モデルによる予備的検証であり、Quantum Hyper-Aetherアーキテクチャ全体の実証を意味するものではない。

Sensitivity Analysis for Finite-VM Knowledge Lifetime Management

Quantum Hyper-Aether simulation summary



EXPERIMENT SCOPE

- Synthetic discrete-event simulation
- Focus: finite VM resources + knowledge lifetime management
- Main outputs: task success rate, retained incorrect knowledge, forced assignment rate
- Sensitivity dimensions:
 - VM capacity $K_{max} = 5, 10, 20, 40$
 - Knowledge TTL = 100, 200, 700, 1500, 3000
 - Verification detection rate = 60%, 75%, 88%, 95%

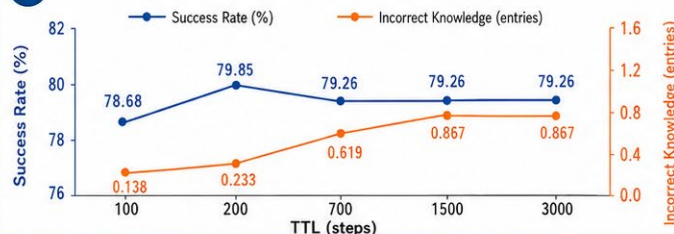
② K_{max} SENSITIVITY

K_{max}	Condition	Success Rate (%)	Incorrect Knowledge	Forced Assignment (%)
5	A: No management	82.86	1.610	15.75
	D: TTL + Verification	82.60	0.025	12.92
10	A: No management	81.37	3.470	12.48
	D: TTL + Verification	81.08	0.141	11.53
20	A: No management	79.57	4.828	10.22
	D: TTL + Verification	79.26	0.619	10.29
40	A: No management	79.45	8.292	7.95
	D: TTL + Verification	80.32	0.473	4.55



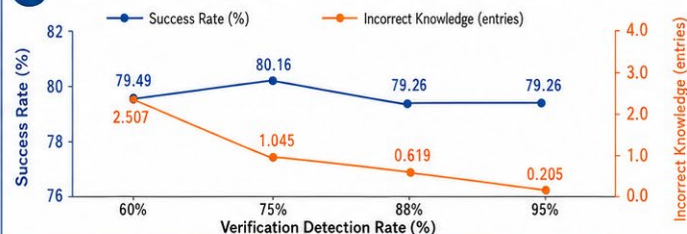
Condition D consistently keeps incorrect knowledge far lower than condition A across all K_{max} .

3A TTL SENSITIVITY UNDER CONDITION D



Short TTL reduces stale incorrect knowledge, but excessively short TTL can slightly lower success. In this setup, TTL $\approx$ 200 gives the best balance.

3B VERIFICATION SENSITIVITY UNDER CONDITION D



Higher verification quality sharply reduces incorrect knowledge. Success does not increase monotonically because strong filtering can also reject some correct knowledge.

KEY FINDINGS



1. Verification is the strongest mechanism for suppressing incorrect knowledge.



2. TTL mainly removes stale knowledge and complements verification.



3. Increasing VM count alone does not guarantee reliability; without management, incorrect knowledge can accumulate.



4. For Quantum Hyper-Aether, VM scaling and knowledge lifetime management should be designed together.



TTL removes old knowledge; verification removes wrong knowledge; together they improve reliability under finite VM resources.

アブストラクト

本研究では、Quantum Hyper-Aetherアーキテクチャにおける有限VM環境を対象として、知識ライフタイム管理の感度分析を行った。シミュレーションでは、有限のVM容量、知識のTTL (Time-to-Live)、および知識検証の相互作用に着目し、合成離散事象モデルを用いて、最大VM数、TTL長、検証による誤知識検出率の変化が、タスク成功率、保持される誤知識量、強制割当率に与える影響を評価した。

結果として、VM数を増やすだけでは信頼性の向上は保証されないことが示された。知識管理を行わない場合、VM容量の増加に伴って保持される知識量が増える一方で、誤知識も蓄積しやすくなる。このことは、実行資源のスケーリングには、知識ライフサイクル制御を組み合わせる必要があることを示している。TTLは主に古い知識を除去する機構として機能し、知識検証は誤知識を直接抑制する機構として機能した。今回の設定では、短いTTLは誤知識の残留を減少させたが、過度に短いTTLでは有用な知識も早期に失効するため、タスク成功率がわずかに低下した。また、検証精度は誤知識抑制に最も強い影響を示し、誤知識検出率が高いほど保持誤知識は大幅に減少した。ただし、強い検証は一部の正しい知識も棄却する可能性があるため、タスク成功率は単調には増加しなかった。

これらの結果は、Quantum Hyper-Aetherにおいて、VMスケーリングと知識ライフタイム管理を一体として設計すべきであることを示している。VMの寿命と知識の寿命を分離し、TTLによる知識失効と知識検証を組み合わせることで、有限な計算資源下でも信頼性を向上させ、誤知識の蓄積を抑制できる可能性がある。ただし、本研究は合成的な意味タスクに基づく予備的検証であり、今後は実際の埋め込みベースの意味タスクおよびLLMを用いた知識検証機構による評価が必要である。

Semantic Routing Simulation Results Summary

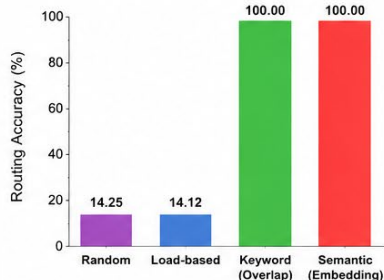
Comparison of Routing Policies on 90 Meaning Tasks Across 8 VM Types (12 Random Seeds)

1. Overall Performance Summary

Routing Policy	Routing Accuracy (%)	Test Set Accuracy (%)	Task Success (%)	Mean Latency (steps)	Mean Semantic Similarity
Random	14.25	14.14	52.79	2.579	0.151
Load-based	14.12	13.82	54.42	2.437	0.150
Keyword (Overlap)	100.00	100.00	85.11	1.880	0.918
Semantic (Embedding)	100.00	100.00	89.13	1.892	0.922

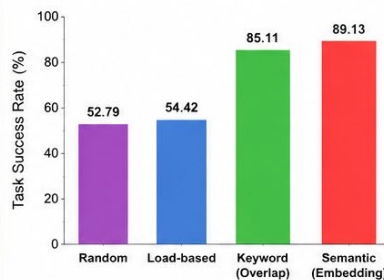
- Semantic routing achieves the highest task success rate (89.13%). It also attains the highest semantic similarity with the expected VM.
- Random and Load-based routing perform close to chance level (~14% routing accuracy) with much lower success rates.
- Keyword routing performs strongly in this dataset, but Semantic routing is slightly better in success and generalization.

2. Routing Accuracy (%)



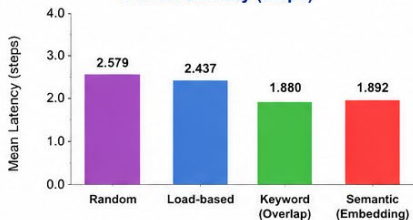
Semantic and Keyword routing correctly send tasks to the expected VM in almost all cases.

3. Task Success Rate (%)



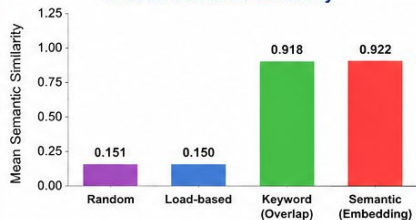
Semantic routing achieves the highest task success rate, demonstrating the benefit of meaning-based routing.

4. Mean Latency (steps)



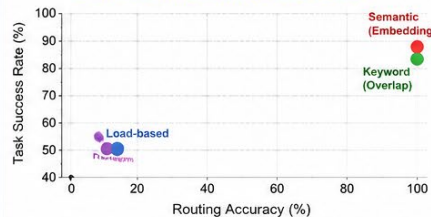
Meaning-based routing (Keyword/Semantic) reduces processing latency compared to non-semantic routing.

5. Mean Semantic Similarity



Semantic routing achieves the highest similarity between task and the selected VM's knowledge.

6. Trade-off: Accuracy vs. Success



Higher routing accuracy leads to higher task success. Semantic routing achieves the best overall trade-off.

Experiment Setup

Dataset: 90 meaning tasks (8 categories)
Train/Test split: 80 / 10

VM Types: 8 specialized VMs

Random Seeds: 12



Routing Policies:
Random, Load-based, Keyword (Overlap),
Semantic (Embedding)



Task Success Model: Based on semantic similarity
and VM specialization



Latency Model: Queueing delay + processing time

Key Takeaways



- Semantic routing outperforms all baselines in task success, latency, and semantic alignment.
- Keyword routing is strong when explicit keywords are available, but semantic routing is more robust to expression variation.
- Load-based routing balances load but ignores meaning, resulting in low accuracy and success.
- These results support the effectiveness of meaning-based routing in Quantum Hyper-Aether.

本研究では、Quantum Hyper-Aetherの意味タスクデータセットを用いて、Random routing、Load-based routing、Keyword-overlap routing、Semantic routingの4つのルーティング方式を比較した。実験では、8種類のVMに対応する90件の意味タスクを使用し、12種類の乱数系列で評価を行った。

結果として、Random routingとLoad-based routingのルーティング精度は約14%にとどまり、タスク成功率はそれぞれ52.79%および54.42%であった。これに対して、Keyword routingとSemantic routingは、このデータセット上でいずれも100.00%のルーティング精度を達成した。Keyword routingのタスク成功率は85.11%であり、Semantic routingは最も高い89.13%のタスク成功率を示した。また、Semantic routingの平均意味類似度は0.922であった。

平均遅延についても、意味に基づく方式では低下が確認された。Keyword routingの平均遅延は1.880ステップ、Semantic routingは1.892ステップであり、Random routingの2.579ステップ、Load-based routingの2.437ステップよりも低かった。

以上の結果から、意味に基づくルーティングは、意味を考慮しない従来型のルーティング方式に比べて、ルーティング適合性とタスク成功率を大きく改善できることが示された。ただし、本実装では、本格的なLLM埋め込みではなく、ドメイン用語とハッシュ特徴に基づく軽量で再現可能な埋め込みを使用している。そのため、今後は実際の埋め込みモデルとLLMベースの検証機構を用いて、より現実的なワークロードで評価する必要がある。

News Semantic Routing Experiment Results Summary

Comparison of Routing Policies on News Semantic Task Dataset (80 Tasks)

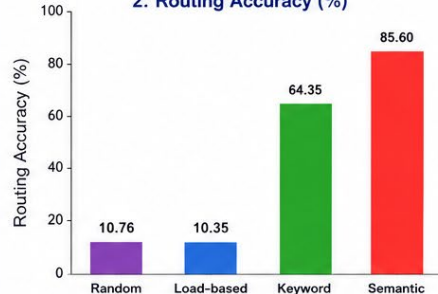
1. Overall Performance Summary

Routing Policy	Routing Accuracy (%)	Test Accuracy (%)	Task Success (%)	Mean Latency (steps)	Mean Semantic Similarity
Random	10.76	10.98	45.68	2.417	0.116
Load-based	10.35	10.23	48.81	2.322	0.112
Keyword	64.35	62.60	64.46	1.977	0.733
Semantic	85.60	77.59	84.01	1.857	0.866

Key Takeaways

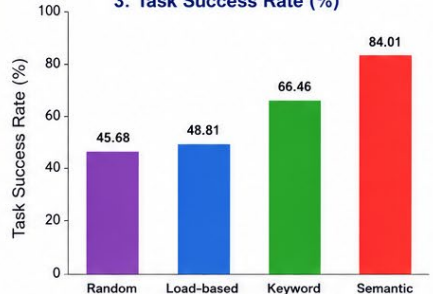
- Semantic routing achieves the best performance across all metrics.
- Keyword routing outperforms Random and Load-based, but is inferior to Semantic.
- Semantic routing improves task success while reducing latency.

2. Routing Accuracy (%)



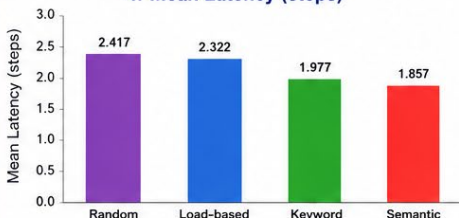
Semantic routing achieves the highest routing accuracy (85.60%), far above Keyword (64.35%) and baseline methods (~10%).

3. Task Success Rate (%)



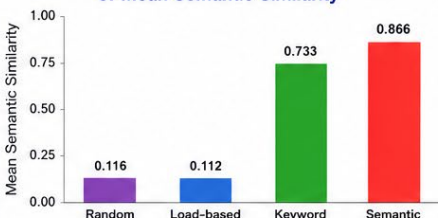
Semantic routing achieves the highest task success rate (84.01%), showing the best ability to complete tasks correctly.

4. Mean Latency (steps)



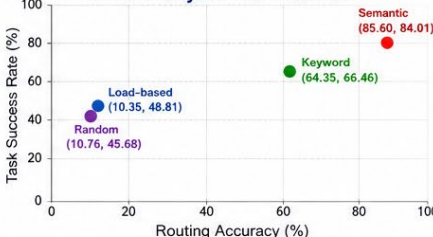
Semantic routing has the lowest latency (1.857 steps), providing more efficient task processing.

5. Mean Semantic Similarity



Semantic routing achieves the highest semantic similarity (0.866), indicating the best semantic alignment with target VMs.

6. Accuracy vs. Task Success



Higher routing accuracy leads to higher task success rate. Semantic routing shows the best overall trade-off.

Experiment Setup

- Dataset: News Semantic Task Dataset (80 tasks)
- VM Profiles: 10 specialized VMs
- Test Set: 16 tasks (20% of total)
- Random Seeds: 12
- Latency Model: Queueing delay + processing time
- Success Model: Based on semantic similarity and VM specialization

VM Profile Categories (10)

- World Politics
- Geopolitical Risk
- Technology & AI
- Economy & Markets
- Security & Public Safety
- Sports & Events
- Climate & Energy
- Health & Science
- Media & Journalism
- Fact-Checking & Verification

Conclusion

Semantic routing significantly outperforms Random, Load-based, and Keyword routing in routing accuracy, task success rate, latency, and semantic alignment on the News Semantic Task Dataset.

These results validate the effectiveness of meaning-based routing for complex, real-world news tasks.

本研究では、Quantum Hyper-Aether向けに設計した **News Semantic Task Dataset** を用いて、Random routing、Load-based routing、Keyword-overlap routing、Semantic routing の4種類のルーティング方式を評価した。このデータセットは、世界政治、地政学リスク、技術・AI、経済・市場、公共安全、気候・エネルギー、健康・科学、メディア・ジャーナリズム、ファクトチェックなど、10種類の専門VMカテゴリに対応する80件のニュース処理タスクから構成される。本実験では、各ルーティング方式がタスクを最適なVMへ割り当てられるかを比較し、ルーティング精度、テスト精度、タスク成功率、平均遅延、意味類似度を測定した。

結果として、Random routingとLoad-based routingはほぼ偶然選択に近い性能にとどまり、ルーティング精度はそれぞれ10.76%および10.35%であった。Keyword routingは、タスク文とVM知識の単語一致を利用することで性能を改善し、64.35%のルーティング精度と66.46%のタスク成功率を達成した。一方、Semantic routingは最も高い性能を示し、ルーティング精度85.60%、テスト精度77.59%、タスク成功率84.01%、平均遅延1.857ステップ、平均意味類似度0.866を達成した。これらの結果は、特に言い換え、曖昧な分類、複数領域にまたがるニュースタスクにおいて、意味類似度に基づくルーティングが、負荷や単純なキーワード一致に基づくルーティングよりも有効であることを示している。

本結果は、Quantum Hyper-Aetherの中核的な考え方、すなわち計算タスクは単にリソースの空き状況だけでなく、タスクの意味と各VMの知識状態との関係に基づいてルーティングされるべきである、という主張を支持するものである。ただし、今回の実装では、本格的な埋め込みモデルではなく、軽量なオントロジーベースの意味ベクトルを使用している。そのため、今後は実際の埋め込みモデルを用いた評価、およびLLMによるファクトチェックや知識検証機構を組み込んだ実験が必要である。

News Semantic Routing Re-evaluation

Comparison of Random, Load-based, Keyword, Ontology Semantic, and Local LSA Semantic routing



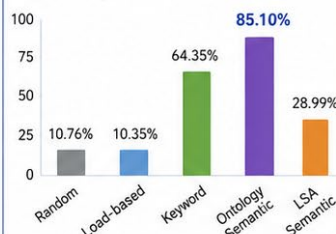
Local runtime note: production embedding models were unavailable, so LSA (TF-IDF + TruncatedSVD) was used as a reproducible local embedding baseline.

Routing Method	Routing Accuracy (%)	Test Accuracy (%)	Task Success (%)	Mean Latency (s)	Semantic Similarity (Mean)
Random	10.76%	10.98%	45.68%	2.417	0.115
Load-based	10.35%	10.23%	48.81%	2.322	0.111
Keyword	64.35%	62.60%	66.46%	1.977	0.732
Ontology Semantic	85.10%	74.40%	83.77%	1.862	0.863
LSA Semantic	28.99%	32.93%	59.66%	2.337	0.606



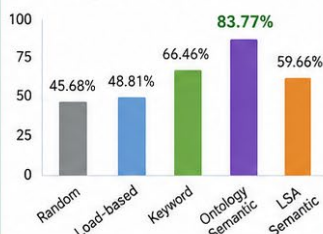
1. Routing Accuracy (%)

Higher is better



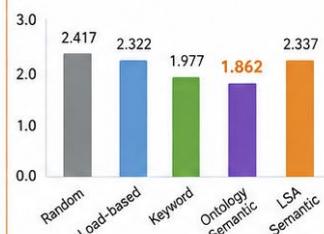
2. Task Success Rate (%)

Higher is better



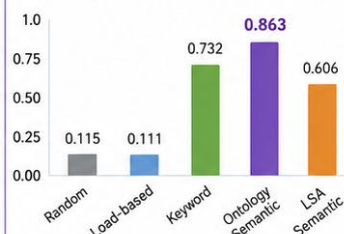
3. Mean Latency (s)

Lower is better



4. Semantic Similarity (Mean)

Higher is better



Key Findings

- Meaning-based routing strongly outperforms non-semantic baselines.
- Ontology Semantic achieved the best overall performance.
- Keyword routing improved performance but remained below ontology-based semantic routing.
- Local LSA embedding improved over Random and Load-based routing, but underperformed due to the small dataset size.



Interpretation / Conclusion

Small datasets are not sufficient for strong statistical semantic embeddings alone. Effective semantic routing in Quantum Hyper-Aether likely requires either stronger production embeddings or ontology-assisted embeddings.



Dataset: 80 news tasks, 10 VM categories, 8 random seeds, 450 simulation steps per seed.

本研究では、Quantum Hyper-AetherにおけるNews Semantic Taskのルーティングを再評価するため、Random、Load-based、Keyword-overlap、Ontology Semantic、Local LSA Semanticの5種類のルーティング方式を比較した。実験では、10種類の専門VMカテゴリに対応する80件のニュース処理タスクを使用し、8種類の乱数系列、各450シミュレーションステップで評価を行った。なお、実行環境では本格的な埋め込みモデルを利用できなかったため、Local LSA Semantic routingでは、TF-IDFとTruncatedSVDを用いた再現可能なローカル埋め込みをベースラインとして実装した。結果として、Random routingとLoad-based routingはほぼ偶然選択に近い性能にとどまり、ルーティング精度はそれぞれ10.76%および10.35%であった。Keyword-overlap routingは性能を大きく改善し、64.35%のルーティング精度と66.46%のタスク成功率を達成した。Ontology Semantic routingは最も高い総合性能を示し、ルーティング精度85.10%、テスト精度74.40%、タスク成功率83.77%、最小平均遅延1.862ステップ、最高平均意味類似度0.863を達成した。一方、Local LSA Semantic routingはRandomおよびLoad-based routingよりは改善したもの、KeywordおよびOntology Semantic routingには及ばず、ルーティング精度28.99%、タスク成功率59.66%であった。これらの結果は、意味に基づくルーティングが、意味を考慮しないベースライン方式と比較して、タスク割当てと処理性能を大きく改善できることを示している。一方で、LSAのような単純な統計的埋め込み手法は、小規模データセットや細かく分かれたニュースカテゴリに対しては不十分である可能性も示された。したがって、Quantum Hyper-Aetherにおける有効なSemantic routingには、より強力な本格的埋め込みモデル、またはオントロジー補助付き意味表現が必要であると考えられる。今後は、実際の埋め込みモデルを用いた同一ベンチマークの評価、およびVM知識更新前のLLMベースのファクトチェックと知識検証機構の統合が必要である。



Revised Definition of Semantic Distance

From embedding-only similarity to ontology-assisted semantic routing

1 Previous Definition

$$D_{sem}(T, VM_i) = D_{emb}(T, VM_i)$$

or $1 - \text{cosine similarity}(\text{task embedding}, \text{VM knowledge embedding})$

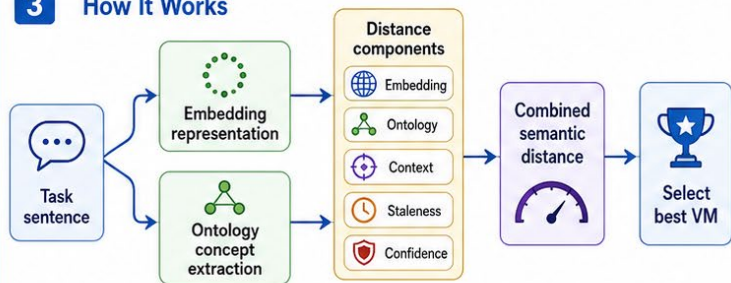
- Embedding-only distance
- Good for simple matching
- Weak on ambiguity, context, and knowledge freshness

2 Extended Definition

$$D_{sem}(T, VM_i) = \alpha D_{emb}(T, VM_i) + \beta D_{ont}(T, VM_i) + \gamma D_{ctx}(T, VM_i) + \delta D_{stale}(T, VM_i) - \varepsilon C_i$$

D_{emb} : embedding distance D_{stale} : knowledge staleness penalty
 D_{ont} : ontology distance C_i : confidence of VM knowledge
 D_{ctx} : context / purpose mismatch $\alpha, \beta, \gamma, \delta, \varepsilon$: tunable weights

3 How It Works



4 Why the New Definition Is Better

- Improves routing for ambiguous and cross-domain tasks
- Adds explainability through ontology structure
- Accounts for knowledge freshness and reliability
- Supports VM specialization more robustly
- Combines statistical meaning and symbolic meaning

Example task: 'Summarize how geopolitical tension could affect energy prices and shipping routes.'



Geopolitical Risk VM:
low distance



Climate / Energy VM:
low distance



Sports VM:
high distance



Routing chooses the VM with the smallest combined semantic distance.

アブストラクト

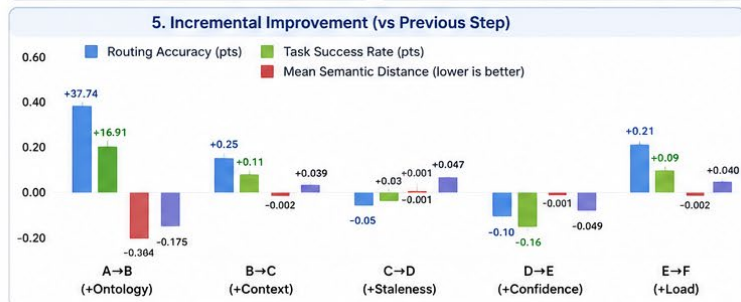
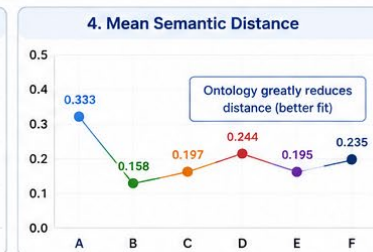
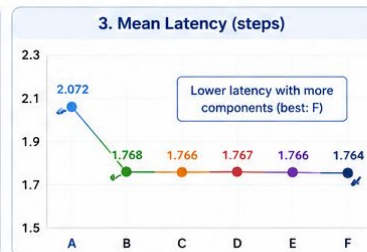
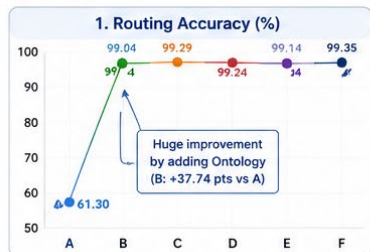
本研究では、Quantum Hyper-Aetherにおけるセマンティック距離の定義を見直し、従来の埋め込みベクトルのみを用いた距離から、オントロジー補助付きの複合距離へと拡張する。従来の定義では、タスクとVMのセマンティック距離は、主にタスク埋め込みとVM知識埋め込みの距離、例えばコサイン類似度の1からの差として計算されていた。この方法は単純な意味マッチングには有効であるが、曖昧なタスク、複数領域にまたがるタスク、文脈依存のタスク、時間変化の影響を受けるタスクには十分ではない。改訂後の定義では、埋め込み距離、オントロジー距離、文脈または目的の不一致、知識の古さ、VM知識の信頼度を組み合わせてセマンティック距離を定義する。オントロジー距離を導入することで、統計的なベクトル類似度に加えて、概念間の構造的関係を利用できるようになる。これにより、Quantum Hyper-Aetherは、地政学リスク、エネルギー市場、公共安全、スポーツイベントのように、互いに関連しながらも異なる領域を、タスク文が間接的または曖昧な場合でも区別できる。さらに、知識の古さと信頼度を距離定義に含めることで、ルーティング判断は単なる意味的近さだけでなく、VMが保持する知識状態の鮮度と信頼性も反映できるようになる。これにより、古い知識や信頼性の低い知識を持つVMへの誤ルーティングを抑制できる。この改訂されたセマンティック距離は、意味ルーティングの説明可能性と頑健性を向上させる。すなわち、システムは単に埋め込みベクトルが近いVMを選ぶのではなく、意味、オントロジー構造、文脈、知識の鮮度、信頼度を考慮した複合的なセマンティック距離が最小となるVMを選択する。したがって、本定義は、Quantum Hyper-AetherにおけるVM専門化、知識ライフタイム管理、信頼性の高いタスクルーティングのための、より強固な基盤を提供する。

Ablation Study of Revised Semantic Distance in Quantum Hyper-Aether

Effect of Each Distance Component on Routing Performance

Ablation Conditions							
ID	Distance Components	Emb.	Ont.	Ctx.	Stale.	Conf.	Load
A	Embedding only	✓	—	—	—	—	—
B	Embedding + Ontology	✓	✓	—	—	—	—
C	Embedding + Ontology + Context	✓	✓	✓	—	—	—
D	Embedding + Ontology + Context + Staleness	✓	✓	✓	✓	—	—
E	Embedding + Ontology + Context + Staleness + Confidence	✓	✓	✓	✓	✓	—
F	Embedding + Ontology + Context + Staleness + Confidence + Load	✓	✓	✓	✓	✓	✓

Overall Results (Averaged over 8 Seeds)					
	Routing Accuracy (%)	Test Accuracy (%)	Task Success Rate (%)	Mean Latency (steps)	Mean Semantic Distance
A	61.30	60.36	79.56	2.072	0.333
B	99.04	98.76	96.47	1.768	0.158
C	99.29	98.76	96.58	1.766	0.197
D	99.24	98.76	96.55	1.767	0.244
E	99.14	98.54	96.39	1.766	0.195
F	99.35	98.84	96.48	1.764	0.235



Key Findings

- Ontology** is the key factor. Adding ontology (A→B) brings the largest improvement in all metrics.
- Context** further improves task success. Helpful for handling task types (verification, hybrid, ambiguous).
- Staleness** and **Confidence** stabilize routing. They prevent outdated or unreliable knowledge from being selected.
- Load** balancing yields the best latency. Condition F achieves the lowest latency and the highest routing accuracy.

Summary

Best Overall Performance
Condition F (all components) achieves the highest routing accuracy (99.35%) and the lowest latency (1.764).

Why It Works
Combining statistical similarity (embedding) with symbolic knowledge (ontology), context, freshness, confidence, and load results in robust routing.

Conclusion
The revised semantic distance definition significantly improves routing accuracy, efficiency, and reliability in Quantum Hyper-Aether.

アブストラクト

本研究では、News Semantic Task Datasetを用いて、Quantum Hyper-Aetherにおける改訂セマンティック距離のアプリケーション分析を実施した。目的は、改訂された距離定義を構成する各要素が、ルーティング性能にどのように寄与するかを評価することである。比較した条件は、(A) Embedding only、(B) Embedding + Ontology、(C) Embedding + Ontology + Context、(D) Embedding + Ontology + Context + Staleness、(E) Embedding + Ontology + Context + Staleness + Confidence、(F) Embedding + Ontology + Context + Staleness + Confidence + Loadの6条件である。実験には、10種類の専門VMカテゴリに対応する80件のニュース処理タスクを使用し、複数の乱数系列で評価を行った。結果として、埋め込みのみを用いたルーティングでは、ルーティング精度61.30%、タスク成功率79.56%、平均遅延2.072、平均セマンティック距離0.333となり、頑健な意味ルーティングには不十分であることが示された。最も大きな改善は、オントロジー情報を追加した場合に得られた。条件Bでは、ルーティング精度が99.04%まで向上し、平均遅延は1.768、平均セマンティック距離は0.158まで低下した。さらに文脈情報を追加した条件Cでは、タスク成功率が96.58%となり、全条件中で最も高い値を示した。すべての要素を含む条件Fでは、ルーティング精度99.35%、テスト精度98.84%、平均遅延1.764となり、総合的に最良の性能を達成した。また、タスク成功率も96.48%と高い水準を維持した。これらの結果は、改訂セマンティック距離において、オントロジーが最も重要な追加要素であることを示している。一方、文脈は曖昧タスクや複合タスクの処理性能を改善し、知識の古さと信頼度は主としてルーティングの信頼性向上に寄与した。負荷情報は、意味的適合性を置き換えるものではないが、運用効率を改善する要素として有効であった。以上より、埋め込み類似度、オントロジー、文脈、知識鮮度、信頼度、負荷を統合した複合セマンティック距離は、埋め込みのみのルーティングと比較して、精度、効率、信頼性の高いVM選択を実現することが示された。本アプリケーション分析は、Quantum Hyper-Aetherにおける意味ルーティングの基盤として、改訂セマンティック距離定義が有効であることを支持する。

Weight Sensitivity Analysis of Revised Semantic Distance

News Semantic Task Dataset

1 Experiment Setup

- 160 weight configurations evaluated
- 4 random seeds per configuration
- 320 simulation steps per seed
- Components: Embedding, Ontology, Context, Confidence, Staleness, Load
- Positive semantic weights normalized; Staleness and Load treated as penalty weights

2 Best Configuration

Best Configuration: W001

Embedding	0.600	Routing Accuracy	Test Accuracy
Ontology	0.300	100.00%	100.00%
Context	0.100		
Confidence	0.000	Task Success Rate	Mean Latency
Staleness	0.000	93.63%	1.754
Load	0.000		

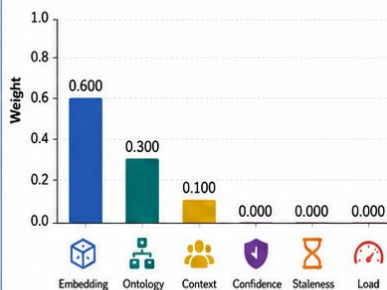
3 Robust High-Performance Region

Robust Region (129 / 160 configurations)
within 0.5 percentage points of the best routing accuracy and task success

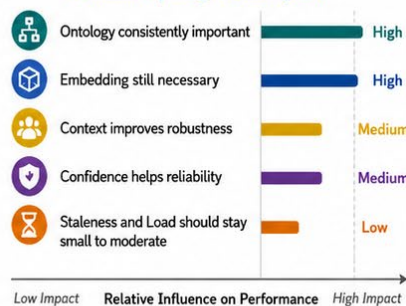
Component	Mean Weight	Min Weight	Max Weight
Embedding	0.327	0.025	
Ontology	0.372	0.181	
Context	0.188	0.016	
Confidence	0.113	0.000	
Staleness	0.103	0.220	
Load	0.101	0.099	

4 Visual Findings

1) Best Configuration Weights (W001)



2) Sensitivity Insights (Conceptual)



5 Key Conclusions

- High performance is not limited to a single weight point.
- Ontology is the most important component for robust routing.
- Balanced Embedding + Ontology + Context is the most reliable design.
- Small penalty weights for Staleness and Load are useful, but should not dominate semantic suitability.

Revised semantic distance is robust across a broad weight region, with ontology-assisted routing as the central driver of performance.

アブストラクト

本研究では、News Semantic Task Datasetを用いて、Quantum Hyper-Aetherにおける改訂セマンティック距離定義の重み感度分析を行った。改訂セマンティック距離は、埋め込み類似度、オントロジー類似度、文脈一致度、VM知識の信頼度、知識の古さ、負荷ペナルティの6成分を統合する。本実験では、改訂距離が特定の手動調整された重み設定に過度に依存していないかを評価するため、160種類の重み設定を比較し、各設定について4種類の乱数系列、320シミュレーションステップで評価した。

最良の設定では、埋め込みに0.600、オントロジーに0.300、文脈に0.100の重みを割り当て、信頼度、知識の古さ、負荷ペナルティは0とした。この設定では、ルーティング精度100.00%、テスト精度100.00%、タスク成功率93.63%、平均遅延1.754を達成した。一方で、高性能は単一の重み設定に限定されなかった。160設定中129設定が、最良のルーティング精度およびタスク成功率から0.5ポイント以内に収まり、広い頑健な重み領域が存在することが示された。

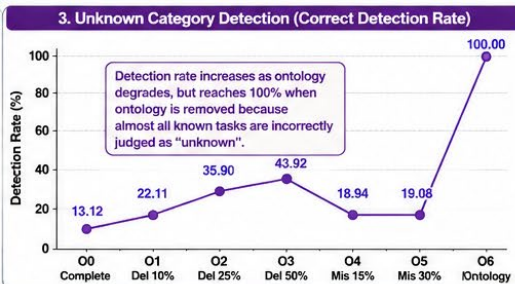
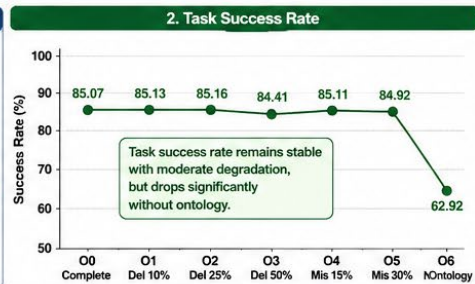
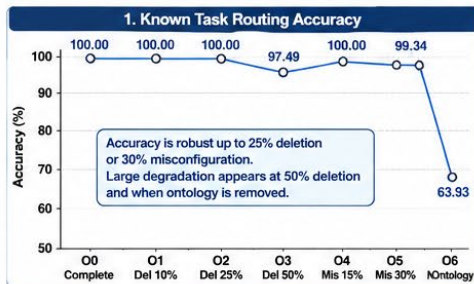
この頑健領域において、オントロジー重みは一貫して重要であり、平均値は0.372、最小値は0.181であった。埋め込みも必要であるが、許容範囲は比較的広く、頑健な意味ルーティングには、統計的類似度とオントロジーに基づく概念構造のバランスが重要であることが分かった。また、文脈成分は、曖昧タスクや複合タスクに対する頑健性を高める役割を果たした。信頼度、知識の古さ、負荷ペナルティは、今回の短時間かつ安定したワークロードでは影響が小さかったが、長時間運用や知識劣化、資源競合が発生する環境では重要になると考えられる。

以上の結果から、改訂セマンティック距離は、単一の手動設定に過度に依存せず、広いパラメータ領域で高い性能を維持できることが示された。特に、オントロジー補助付きルーティングが頑健性の中心的要因である。本研究の結果は、複合セマンティック距離が、Quantum Hyper-Aetherにおける信頼性の高いVM選択の安定した基盤となることを支持している。

Ontology Robustness Experiment in Quantum Hyper-Aether

Evaluation of routing performance under ontology degradation: link deletion, misconfiguration, and unknown category addition

Condition (Topology State)	Description	Known Tasks (Existing 10 Categories)			Unknown Category Detection (8 disaster response tasks)		Mean Latency (lower is better)
		Routing Accuracy	Test Accuracy	Task Success Rate	Detection Rate (Correct)	False Unknown Rate (Known tasks)	
✓ O0 Complete Ontology	All concept and category links are correct	100.00%	100.00%	85.07%	13.12%	0.00%	1.758
10% O1 Delete 10% Links	Randomly delete 10% of concept links	100.00%	100.00%	85.13%	22.11%	0.00%	1.758
25% O2 Delete 25% Links	Randomly delete 25% of concept links	100.00%	100.00%	85.16%	35.90%	0.00%	1.756
50% O3 Delete 50% Links	Randomly delete 50% of concept links	97.49%	96.92%	84.41%	43.92%	2.02%	1.770
✗ O4 Misconfigure 15% Links	Randomly misconfigure 15% of concept links	100.00%	100.00%	85.11%	18.94%	0.13%	1.755
✗ O5 Misconfigure 30% Links	Randomly misconfigure 30% of concept links	99.34%	99.63%	84.92%	19.08%	0.97%	1.758
✗ O6 No Ontology	Remove entire ontology (embedding + context only)	63.93%	64.04%	62.92%	100.00%	99.97%	1.981



Key Findings

- ✓ The revised semantic distance is robust to moderate ontology degradation.
- ✓ Up to 25% link deletion or 30% misconfiguration, performance remains almost unchanged.
- ✓ At 50% link deletion, accuracy drops by about 2.5 points, but task success is mostly preserved.
- ✓ Removing the ontology causes a drastic performance collapse (≈64% accuracy).
- ✓ Unknown-category detection needs a dedicated mechanism (Unknown / Discovery VM + score margin + confidence + LLM verification).

Summary



Takeaway

Ontology-assisted routing is the core driver of high performance in Quantum Hyper-Aether. The system tolerates moderate ontology errors, but requires confidence-aware scoring and an Unknown / Discovery VM to handle severe degradation and truly novel topics.

Experiment Settings: News Semantic Task Dataset (80 existing tasks + 8 unknown tasks)
4 random seeds × 320 steps per run



Robust



Degradation begins



Severe degradation

Lower latency is better.

アブストラクト

本研究では、Quantum Hyper-Aetherにおけるオントロジー補助付きセマンティックルーティングの頑健性を、オントロジー劣化条件下で評価した。実験には、80件の既存ニュース処理タスクからなるNews Semantic Task Datasetに加え、未知カテゴリであるdisaster_responseに属する8件の追加タスクを使用した。比較条件として、完全なオントロジー、概念リンクの10%、25%、50%削除、概念リンクの15%、30%誤設定、およびオントロジーを完全に除去した条件を設定した。評価指標には、既知タスクのルーティング精度、テスト精度、タスク成功率、未知カテゴリ検出率、誤Unknown率、平均遅延を用いた。

結果として、改訂セマンティック距離は中程度のオントロジー劣化に対して高い頑健性を示した。完全なオントロジーでは、既知タスク精度とテスト精度はいずれも100.00%、タスク成功率は85.07%であった。概念リンクを10%または25%削除しても、既知タスク精度とテスト精度は100.00%を維持し、タスク成功率も約85%で安定した。また、概念リンクを30%誤設定した場合でも、既知タスク精度は99.34%、テスト精度は99.63%を維持した。一方、概念リンクを50%削除すると、既知タスク精度は97.49%、テスト精度は96.92%に低下し、明確な性能劣化が始まった。さらに、オントロジーを完全に除去すると、既知タスク精度は63.93%、テスト精度は64.04%、タスク成功率は62.92%まで大きく低下した。未知カテゴリ検出では、単純なしきい値による判定だけでは不十分であることが示された。オントロジー劣化が大きくなるほど未知カテゴリ検出率は上昇し、オントロジーなしでは100.00%に達した。しかし同時に、既知タスクの99.97%が誤ってUnknownと判定されるという問題が発生した。この結果は、未知カテゴリ処理には、単なるスコアしきい値ではなく、専用のUnknown / Discovery VM、候補間マージン、オントロジー信頼度、LLM検証を組み合わせた仕組みが必要であることを示している。

以上より、オントロジー補助は高精度な意味ルーティングの中心的要素であり、EmbeddingとContextは中程度のオントロジー損傷を補完できることが確認された。一方で、重度の劣化や新規トピックに対応するためには、オントロジーリンク信頼度、自動修復機構、Unknown / Discovery VMをQuantum Hyper-Aetherに追加する必要がある。



Unknown / Discovery VM Comparison in Quantum Hyper-Aether

Corrected evaluation of unknown-category detection and handling



1 Overview



Dataset: News Semantic Task Dataset

Known classes: 10 VM categories

Unknown class: disaster_response

Compared methods: 6

Key correction: after a new VM is created, unknown tasks handled directly by the new VM are counted in Total Unknown Handling Rate.

2 Comparison of Unknown-Handling Methods



Condition	Known Accuracy	Unknown Precision	Total Unknown Handling	Integrated Unknown F1	False Unknown	Task Success	Mean Latency (s)
A No unknown mechanism	100.00%	0.00%	0.00%	0.00%	0.00%	87.97%	1.770
B Top-1 threshold	100.00%	0.00%	0.00%	0.00%	0.00%	87.97%	1.770
C + Margin	100.00%	100.00%	11.47%	20.58%	0.00%	88.63%	1.800
D + Ontology confidence	99.90%	96.30%	12.80%	22.59%	0.10%	88.71%	1.803
E + Verification	99.74%	93.58%	17.12%	28.95%	0.26%	88.94%	1.815
F Discovery + new VM	99.76%	78.19%	79.88%	79.02%	0.24%	93.24%	1.786

4 Key Findings



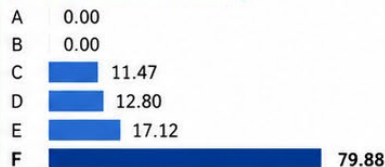
- Simple threshold rules did not handle unknown tasks effectively.
- Margin, ontology confidence, and verification improved unknown detection step by step.
- Discovery VM plus automatic new-VM creation achieved the best overall result.
- The best method preserved known-task accuracy at 99.76% while raising total unknown handling to 79.88%.
- Overall task success increased to 93.24%, showing that Quantum Hyper-Aether can extend its semantic structure to new categories.

3 Visual Summary

1. Integrated Unknown F1



2. Total Unknown Handling



3. Task Success



Conclusion:
Unknown / Discovery VM is essential for scalable semantic expansion.

アブストラクト

本研究では、Quantum Hyper-Aetherにおける未知カテゴリの検出および処理能力を評価するため、6つのルーティング条件を比較した。比較条件は、

(A) 未知カテゴリ処理機構なし、(B) 第1候補スコアのしきい値判定、(C) しきい値と第1・第2候補間マージンの併用、(D) マージンとオントロジー信頼度の併用、(E) オントロジー信頼度と検証機構の併用、(F) Unknown / Discovery VMと新規VMの自動生成を組み合わせた方式である。実験には、10種類の既知VMカテゴリから構成される News Semantic Task Datasetと、未知カテゴリとして追加した disaster_response タスクを使用した。

新規生成されたVMは、それ以降の未知タスクを Discovery VMを介さず直接処理できるため、評価には Precision や Recall に加えて、補正指標である Total Unknown Handling Rate (未知タスク総処理率) を導入した。結果として、単純なしきい値判定だけでは未知カテゴリを十分に処理できない一方、第1・第2候補間マージン、オントロジー信頼度、検証機構を段階的に追加することで、未知カテゴリ検出能力が改善することが確認された。

最も高い性能を示したのは、Discovery VMと新規VM生成を組み合わせた条件Fであった。この方式は、既知タスクのルーティング精度99.76%、未知タスク総処理率79.88%、統合Unknown F1スコア79.02%、誤Unknown率0.24%、タスク成功率93.24%を達成した。

以上の結果は、Quantum Hyper-Aetherが既知タスクに対する高い処理性能を維持しながら、Discovery VMによる未知意味領域の検出と、専門VMの生成によって、新しいセマンティック領域をシステム内部へ取り込める可能性を示している。

Lifecycle Management of Provisional VMs in Quantum Hyper-Aether

Summary of Experimental Results

Condition	Lifecycle Strategy	Lifecycle Action Accuracy (%) ↑	Promote Accuracy (%) ↑	Merge Accuracy (%) ↑	Reject Accuracy (%) ↑	Delete Accuracy (%) ↑	Wrong Promotion Rate (%) ↓	# of Permanent VMs (avg) ↓	# of Duplicate VMs (avg) ↓	Task Success Rate (%) ↑
A	Immediate VM creation (at first detection)	20.06	20.06	–	–	–	79.94	88.50	74.50	63.99
B	Promote by count only	26.36	26.36	–	–	–	73.64	47.00	33.25	72.09
C	Count + Cluster cohesion	27.15	27.15	–	–	–	72.85	32.12	19.88	74.23
D	+ Verification	40.47	40.47	–	–	–	59.53	19.38	11.25	74.40
E	+ Merge / Reject	98.51	95.09	100.00	100.00	–	4.91	5.12	2.38	76.47
F	Full lifecycle (E + Delete)	64.16	95.09	100.00	100.00	57.63	4.91	5.12	2.38	76.47

↑ Higher is better. ↓ Lower is better.

Experiment Overview



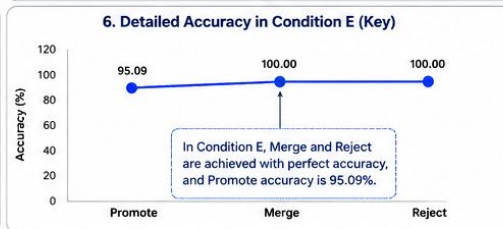
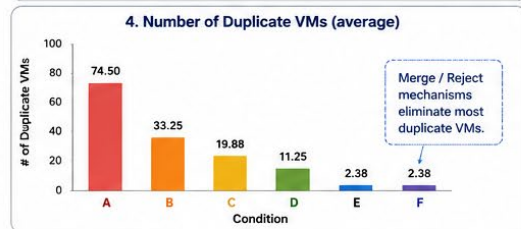
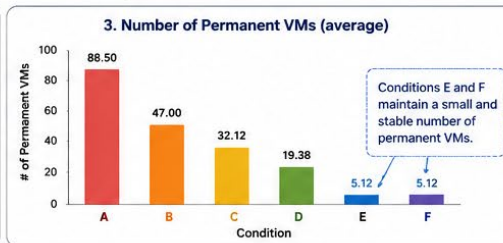
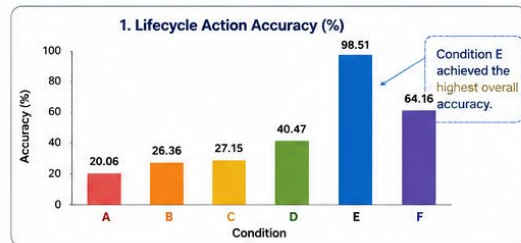
- Dataset
 - News Semantic Task Dataset
 - Mixture of tasks:
 - New categories to promote: disaster_response, transportation_disruption, legal_regulation
 - Tasks to merge: similar to existing VMs (e.g., Technology & AI, Climate & Energy)
 - Noise tasks to reject
 - Regular known news tasks



Number of Runs: 30 per condition



Key Goal: Manage provisional VMs effectively (Promote / Merge / Reject / Delete)



Key Findings

- Immediate creation (A) causes an explosion of VMs and duplicates.
- Count and cohesion alone (B, C) are insufficient to prevent wrong promotions.
- Verification (D) improves performance, but is not enough.
- Merge and Reject (E) are the most critical for high accuracy and low duplication.
- Delete (F) is necessary but needs more intelligent criteria (beyond idle time) to avoid deleting potentially valuable clusters.

Recommendation: Use Condition E (Merge + Reject with Verification and Similarity Check) as the core lifecycle policy.

Employ a confidence-aware delete strategy to safely remove truly obsolete provisional VMs while preserving future-value clusters.

アブストラクト

本研究では、Quantum Hyper-Aetherにおいて、Discovery VMが生成した仮想マシンのライフサイクル管理方式を評価した。比較した方式は、(A) 仮VMの即時生成、(B) タスク数のみに基づく昇格、(C) タスク数とクラスタ凝集度に基づく昇格、(D) 検証機構を追加した昇格、(E) 検証に加えてMergeおよびReject判定を行う方式、(F) さらにDeleteを含む完全なライフサイクル方式の6条件である。実験ワークロードには、通常の既知ニュースタスクに加え、真に新しい意味カテゴリである disaster_response、

transportation_disruption、legal_regulation、既存VMへ統合すべきタスク、およびRejectすべき無関係なノイズタスクを含めた。結果として、仮VMの即時生成は深刻なVM増殖を引き起こした。条件Aでは、平均88.50個の正式VMが生成され、そのうち74.50個が重複VMであり、誤昇格率は79.94%に達した。タスク数のみ、またはタスク数とクラスタ凝集度に基づく条件BおよびCでは、VM総数は減少したものの、誤昇格率は依然として72%を超えた。検証を追加した条件Dでは、昇格精度が40.47%まで改善したが、信頼できるライフサイクル制御としては不十分であった。

最も高い総合性能を示したのは、検証、既存VMとの重複判定に基づくMerge、およびノイズのRejectを組み合わせた条件Eであった。条件Eは、ライフサイクル判断精度98.51%、Promote精度95.09%、Merge精度100.00%、Reject精度100.00%を達成した。誤Promote率は4.91%まで低下し、正式VM数と重複VM数はそれぞれ平均5.12個、2.38個に抑えられた。また、タスク成功率も全条件中で最も高い76.47%となった。

完全なライフサイクル方式である条件Fは、Promote、Merge、Rejectについては条件Eと同じ性能を維持した。しかし、単純なDelete規則の精度が57.63%にとどまり、将来利用価値を持つ可能性のある仮VMクラスタまで削除したため、ライフサイクル全体の判断精度は64.16%へ低下した。

以上の結果から、タスク数や意味クラスタの凝集度だけでは、仮VMの増殖を十分に制御できないことが示された。信頼性の高いライフサイクル管理には、正式昇格前の検証、既存VMとの意味的重複判定、明示的なMergeおよびReject処理が必要である。Delete機構も長期的な資源管理には必要であるが、単純な未使用期間だけではなく、信頼度、再利用可能性、新規性、将来価値を考慮する必要がある。本研究の結果は、Quantum Hyper-Aetherにおける拡張可能な意味構造管理の中核方式として、検証付きPromote、Merge、Rejectを採用することの有効性を支持する。

Improved Delete Mechanism Experiment

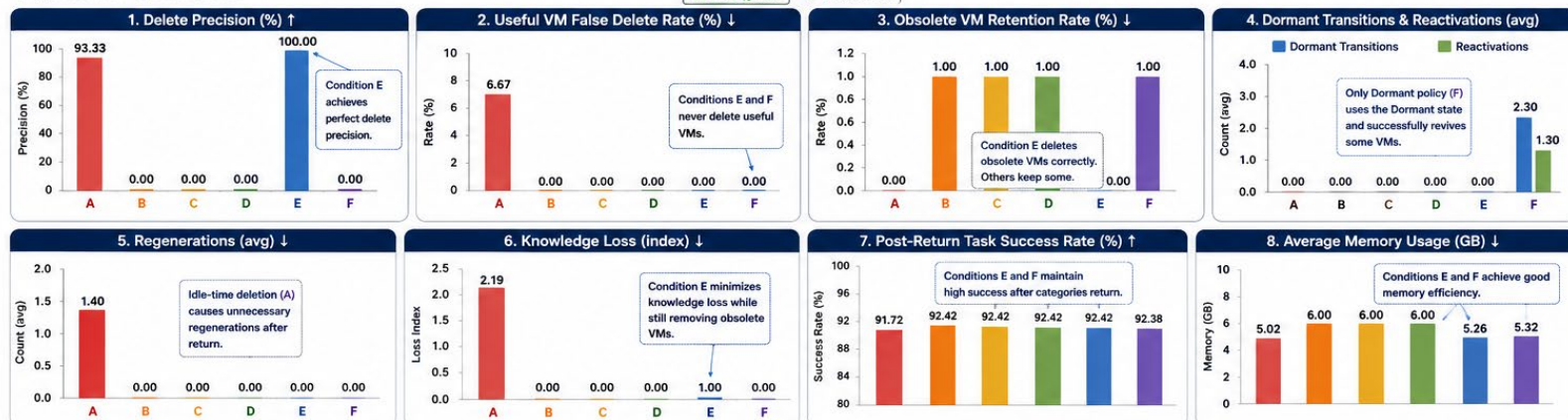
Dormant State, Semantic Novelty, and Reuse Probability

SUMMARY OF RESULTS

Condition (Delete Policy)	Delete Precision (%) ↑	Useful VM False Delete Rate (%) ↓	Obsolete VM Retention Rate (%) ↓	Dormant Transitions (avg)	Reactivations from Dormant (avg)	Regenerations (avg) ↓	Knowledge Loss (index) ↓	Post-Return Task Success Rate (%) ↑	Average Memory Usage (GB) ↓
A Idle-time only	93.33	6.67	0.00	0.00	0.00	1.40	2.19	91.72	5.02
B Idle-time + Usage	0.00	0.00	1.00	0.00	0.00	0.00	0.00	92.42	6.00
C + Confidence	0.00	0.00	1.00	0.00	0.00	0.00	0.00	92.42	6.00
D + Semantic Novelty	0.00	0.00	1.00	0.00	0.00	0.00	0.00	92.42	6.00
E + Reuse Probability	100.00	0.00	0.00	0.00	0.00	0.00	1.00	92.42	5.26
F Dormant State + Confidence-aware Delete	0.00	0.00	1.00	2.30	1.30	0.00	0.00	92.38	5.32

↑ Higher is better. ↓ Lower is better.

Best (good) Worst (bad)



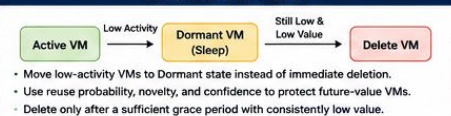
KEY FINDINGS

- ✓ Idle-time only deletion (A) removes some useful VMs, causing regenerations, knowledge loss, and lower success after return.
- ✓ Usage, confidence, and novelty alone (B-D) are too conservative and fail to delete obsolete VMs.
- ✓ Adding reuse probability (E) achieves the best balance: perfect delete precision, no false deletions, no regenerations, low memory usage, and high post-return success.
- ✓ Dormant state (F) protects valuable VMs through a sleep-wake mechanism, eliminating knowledge loss and regenerations.

DELETE DECISION FACTORS (Relative Importance)



RECOMMENDED POLICY



The combination of Reuse Probability and Dormant State enables safe deletion of unnecessary VMs while preserving future-value knowledge.

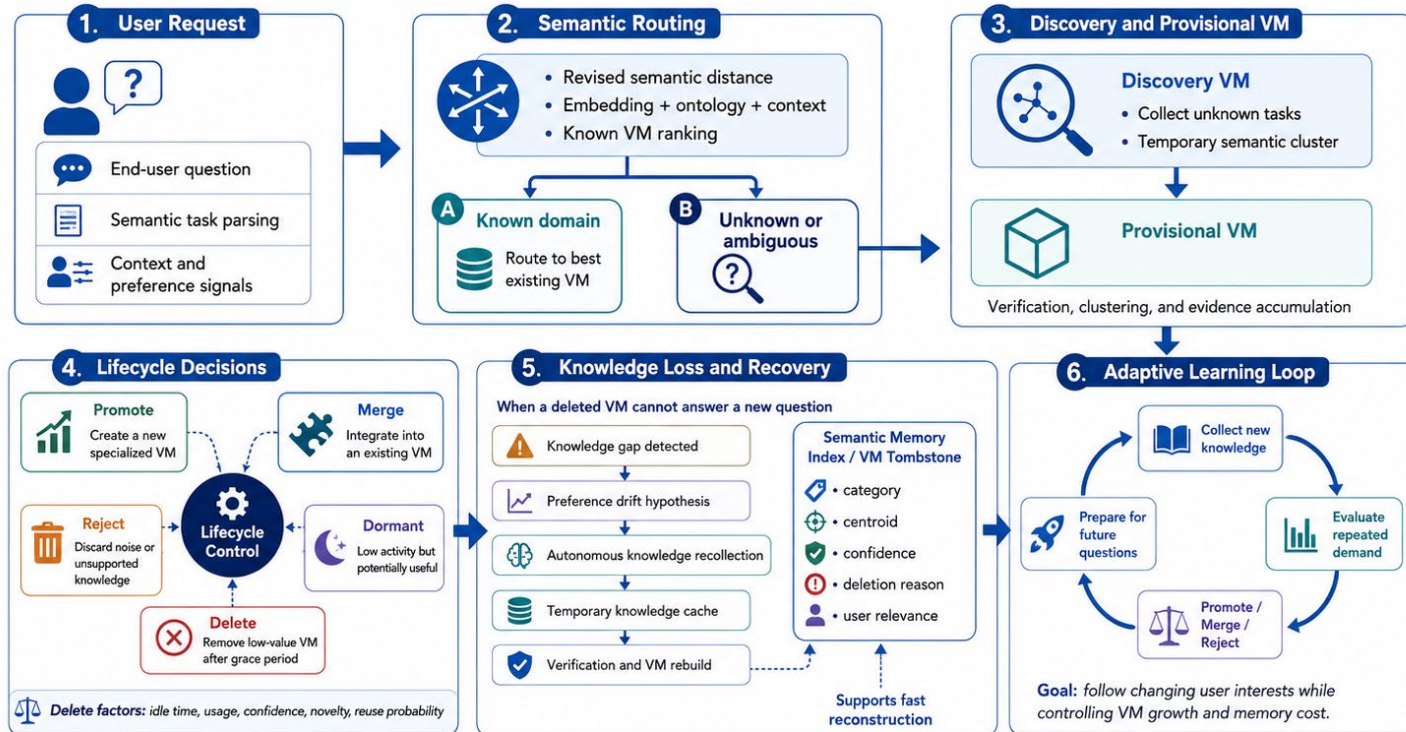
アブストラクト

本研究では、Quantum Hyper-AetherにおけるセマンティックVMの改良Delete機構を評価し、Dormant状態、意味的新規性、再利用確率がVM削除判断に与える効果を検証した。実験では、ある意味カテゴリが一時的に消滅した後に再び出現し、別のカテゴリは完全に不要になるという時間変化を含むワークロードを用いた。比較した方式は、(A) 未使用時間のみに基づく削除、(B) 未使用時間と使用回数、(C) 使用回数と信頼度、(D) 信頼度と意味的新規性、(E) 新規性と再利用確率、(F) Dormant状態と信頼度を考慮した削除、の6条件である。結果として、単純な未使用時間に基づく削除は危険であることが示された。条件Aは93.33%のDelete精度を達成したが、有用VMを6.67%の割合で誤削除し、その結果、VMの再生成、知識損失、カテゴリ復帰時の頑健性低下を引き起こした。条件B、C、Dは保守的すぎる挙動を示し、有用VMの誤削除は避けられたものの、不要になったVMを削除できず、余分なメモリ使用を維持した。最も良い直接削除性能を示したのは、再利用確率を導入した条件Eであった。条件Eは、Delete精度100.00%、有用VMの誤削除率0.00%、VM再生成0回、低いメモリ使用量、復帰後タスク成功率92.42%を達成した。一方、Dormant状態を導入した条件Fは、知識損失を0に抑え、VM再生成なしに再活性化を可能にした。しかし、今回の保守的な設定では、不要VMを削除するよりも保持する傾向が強く、完全に消滅したVMが残存する課題も確認された。

これらの結果から、VMを削除すべきかどうかを判断するうえで、再利用確率が最も有効な要素であることが示された。一方で、Dormant状態は、出現頻度は低いが将来再利用される可能性があるVMを保護する中間状態として有効である。したがって、実用上の最適な方針は、低活動VMをまずDormant状態へ移行し、その後、十分な猶予期間を経ても再利用確率、信頼度、将来価値が低い場合にのみ削除する方式である。本結果は、Quantum Hyper-Aetherにおける意味的リソース管理において、より安全で適応的なVMライフサイクル制御が必要であることを示している。

VM Operations in Quantum Hyper-Aether

Semantic routing, lifecycle control, knowledge recovery, and adaptive learning



Operational Principles



Meaning-based routing
Route by semantic meaning and context, not keywords.



Controlled VM growth
Create only when valuable; merge and retire responsibly.



Safe forgetting with recovery
Delete to save resources, recover when knowledge is needed.



Adaptive user-centered learning
Evolve with user interests and real-world feedback.

アブストラクト

この図は、Quantum Hyper-AetherにおけるVM運用を、エンドツーエンドの意味的ライフサイクルとしてまとめたものである。エンドユーザの要求は、まずセマンティックタスクとして解析され、文脈情報や嗜好信号とともに整理される。その後、埋め込み、オントロジー、文脈情報を組み合わせた改訂セマンティック距離に基づいてルーティングされる。タスクが既知領域に属する場合は、最も適した既存VMへ送られる。一方、未知または曖昧なタスクの場合はDiscovery VMへ送られ、検証と証拠蓄積を通じて、一時的な意味クラスタや仮VMが形成される。

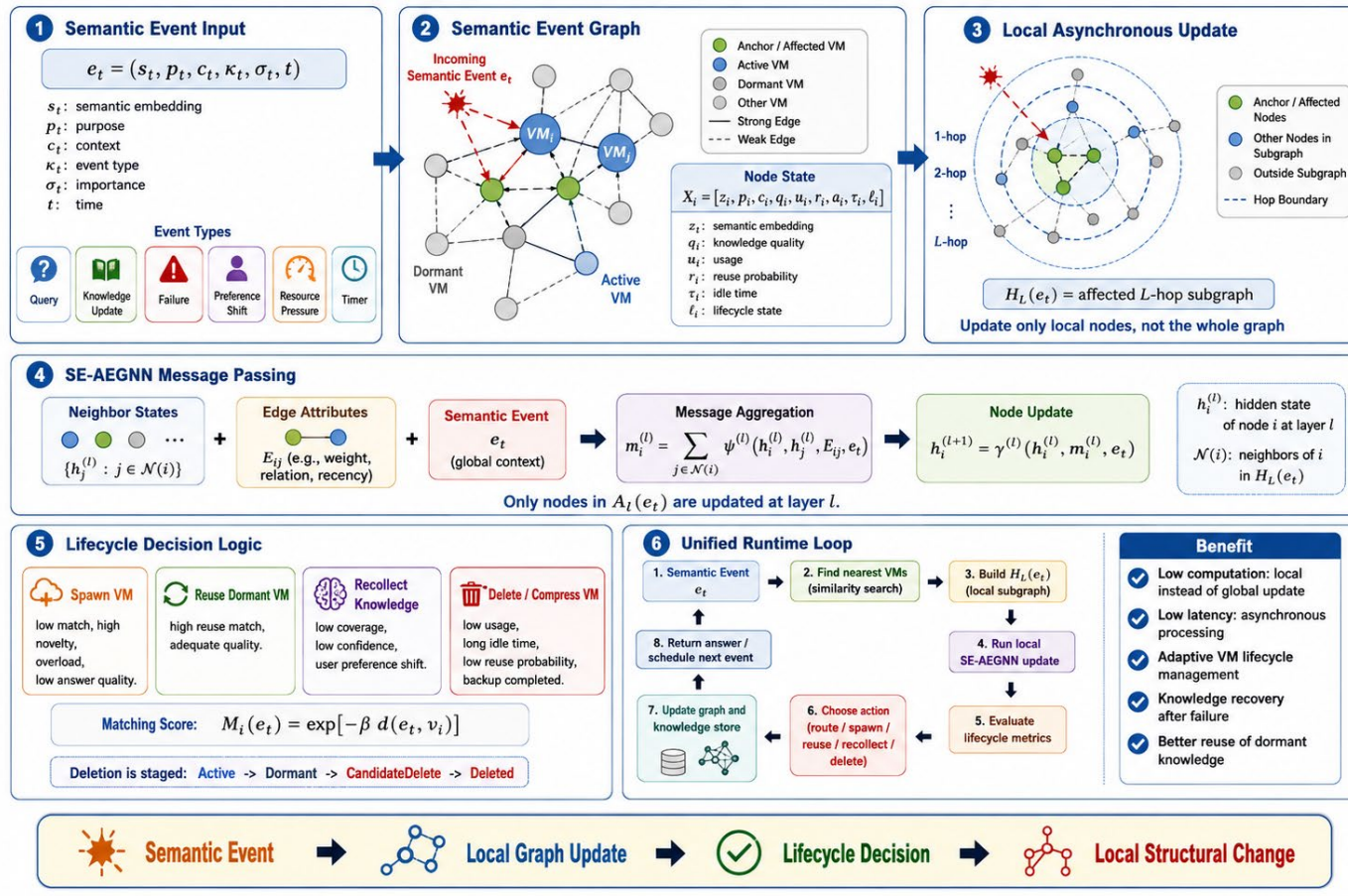
VMのライフサイクルは、Promote、Merge、Reject、Dormant、Deleteの5つの判断によって制御される。Promoteは安定した新規領域に対して専門VMを生成する処理であり、Mergeは既存VMと重複する知識を統合する処理である。Rejectは根拠のない知識やノイズを破棄し、Dormantは活動頻度は低い将来有用となる可能性のあるVMを休眠状態として保持する。Deleteは、未使用時間、使用頻度、信頼度、新規性、再利用確率などに基づき、十分な猶予期間後に価値の低いVMを削除する処理である。

VM削除によって知識ギャップが発生し、その後のユーザ要求に回答できなくなった場合、システムは知識不足を検出し、ユーザの嗜好や関心が変化した可能性を仮説として扱う。そのうえで、自律的な知識再収集を開始する。削除されたVMについては、カテゴリ、意味重心、信頼度、削除理由、ユーザ関連度などを保持する軽量のSemantic Memory IndexまたはVM Tombstoneを残すことで、必要に応じて高速に再構築できるようにする。再収集された知識は一時キャッシュに保存され、検証後に適切なVMの再構築または更新に用いられる。

この一連の処理により、システムは新しい知識を収集し、繰り返し需要を評価し、Promote、Merge、Rejectを通じてVM構造を更新する適応学習ループを形成する。最終的な目的は、意味に基づくルーティング、制御されたVM成長、安全な忘却と回復、そしてユーザ中心の適応学習を実現しながら、意味的精度、資源効率、長期的な知識進化のバランスを取ることである。

Semantic Event-based Asynchronous Graph Neural Network (SE-AEGNN) for Quantum Hyper-Aether

Local graph updates for VM lifecycle management and knowledge adaptation



要旨

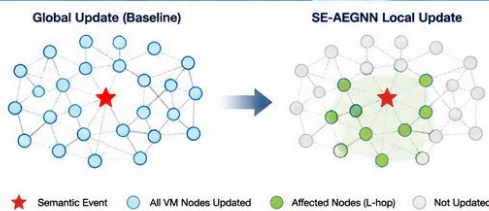
本研究では、Quantum Hyper-AetherにおけるVMライフサイクル管理と知識適応を実現するため、Semantic Event-based Asynchronous Graph Neural Network (SE-AEGNN)を提案する。本方式では、入力される意味イベントを、意味埋め込み、目的、文脈、イベント種別、重要度、時刻から構成される構造化データとして表現する。これらのイベントは、VMをノード、意味的關係、再利用履歴、文脈的関連性をエッジとして持つ意味イベントグラフへ写像される。SE-AEGNNはグラフ全体を再計算するのではなく、イベントの影響を受ける局所的な L -hop部分グラフのみを更新することで、低遅延かつ効率的な非同期処理を実現する。メッセージ伝播では、近傍VMの状態、エッジ属性、現在の意味イベントを統合し、局所的に関連するノードだけを更新する。更新結果に基づき、システムはVM生成、休眠VMの再利用、知識再収集、段階的なVM削除・圧縮を実行する。さらに、類似度探索、局所部分グラフ生成、非同期更新、ライフサイクル評価、グラフおよび知識ストアの更新を統合した実行ループを構成する。これにより、SE-AEGNNは、全体更新に比べて計算量を抑えながら、適応的な意味スケジューリング、回答失敗後の知識回復、休眠知識の効率的再利用を実現する。

REF. AEGNN: Asynchronous Event-based Graph Neural Networks Simon Schaefer*

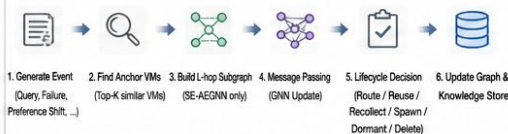
SE-AEGNN vs. Global GNN Update: Simulation Results Summary

SE-AEGNN (Semantic Event-based Asynchronous Graph Neural Network) updates only the affected L-hop subgraph, achieving the same decisions as Global Update with much lower computation.

1. Concept: Local Update in Semantic Event Graph



2. Simulation Workflow



7. Simulation Settings

- VM counts: 100, 300, 1000
- Embedding dimension: 32
- Average degree: ~8
- SE-AEGNN layers (L): 2
- Anchor VMs (Top-K): 3
- Semantic events: 300
- Event types:
 - Query (55%)
 - Knowledge Update (10%)
 - Failure (10%)
 - Preference Shift (10%)
 - Resource Pressure (10%)
 - Timer (5%)

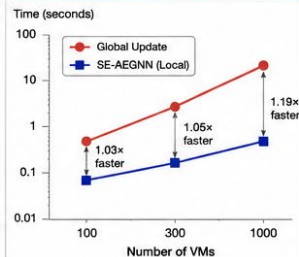
8. Lifecycle Decisions

- Route**: Use existing active VM to answer.
- Reuse**: Reactivate a dormant VM.
- Recollect**: Collect additional knowledge.
- Spawn**: Create a new VM.
- Dormant**: Move VM to dormant state.
- Delete**: Compress and remove VM.

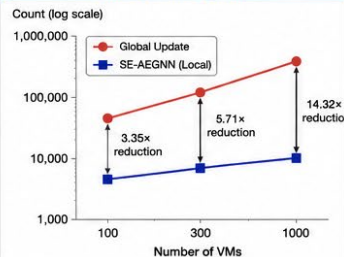
3. Quantitative Comparison (Average over 300 Semantic Events)

# of VMs	Avg. Local Updated VMs (SE-AEGNN)	Local Update Ratio ($ H_L / N$)	Node Evaluation Count (Total)		Evaluation Count Reduction (Global / Local)	Runtime Speedup (Global / Local)	Decision Agreement (with Global)
			Global (Update All)	SE-AEGNN (Local)			
100	89.5	89.5%	60,600	18,100	3.35×	1.03×	100%
300	157.5	52.5%	181,800	31,850	5.71×	1.05×	100%
1000	209.5	20.9%	606,000	42,350	14.32×	1.19×	100%

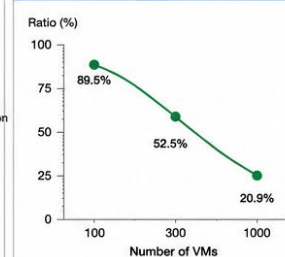
4. Runtime Comparison (Total Time for 300 Events)



5. Node Evaluation Count (Total over All Layers)



6. Local Update Ratio ($|H_L| / N$)

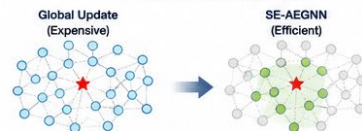


9. Key Findings

- ✓ SE-AEGNN updates only a small portion of the graph.
- ✓ As the number of VMs increases, the local update ratio decreases significantly.
- ✓ SE-AEGNN achieves 3.35x to 14.32x fewer node evaluations.
- ✓ Runtime is consistently reduced (up to 1.19x faster in this minimal Python simulation).
- ✓ Lifecycle decisions are 100% consistent with the Global Update in this experiment.
- ✓ SE-AEGNN enables scalable and efficient VM lifecycle management without sacrificing decision quality.

10. Takeaway

SE-AEGNN provides the same lifecycle decisions as Global Update while dramatically reducing the computational scope.



Note: This is a minimal simulation with random semantic vectors and rules. In practical systems with heavier GNN/LLM computations and larger graphs, the runtime advantage of SE-AEGNN is expected to be much more significant.

要旨

本研究では、Quantum Hyper-Aetherにおける意味イベント処理を対象として、SE-AEGNNによる局所更新方式と、グラフ全体を更新するGlobal GNN方式を比較する最小Pythonシミュレーションを実施した。実験では、イベントの影響を受ける L-hop部分グラフのみに計算を限定することで、処理効率とVMライフサイクル判断の一貫性がどのように変化するかを評価した。シミュレーション条件は、VM数を100、300、1000、意味埋め込み次元を32、グラフ層数を2、意味イベント数を300とした。結果として、SE-AEGNNはGlobal方式と同一のライフサイクル判断を維持しながら、計算対象を大幅に削減した。局所更新されるVMの平均割合は、100 VMで89.5%、300 VMで52.5%、1000 VMで20.9%となり、グラフ規模が大きくなるほど局所性の効果が高まることが示された。また、ノード評価回数はそれぞれ3.35倍、5.71倍、14.32倍削減され、実行時間も最小実装上で1.03倍から1.19倍高速化した。さらに、ライフサイクル判断の一致率は100%、状態誤差は0であり、本実験条件では局所更新がGlobal方式と同一の結果を再現できることが確認された。これらの結果は、SE-AEGNNが判断品質を維持しながら計算量を削減できることを支持しており、GNN、LLM、知識検索などの処理負荷がより大きい実システムでは、さらに大きな性能向上が期待できる。

SE-AEGNN LifeCycle Experiment: 4 Methods × 3 Scenarios Results Summary

SE-AEGNN = Semantic Event-based Asynchronous Graph Neural Network

1. Four Methods Compared

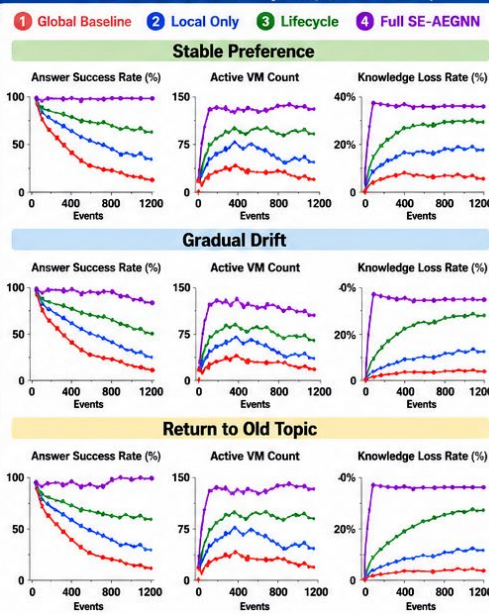
- 1 Global Baseline**
Global update, no reuse, no recollection, no staged delete
- 2 Local Update Only**
Local update only (no reuse, no recollection)
- 3 Lifecycle SE-AEGNN**
Local update + Reuse (Dormant VM) + Staged delete
- 4 Full SE-AEGNN**
Local update + Reuse + Knowledge recollection + Staged delete

2. Average Performance over 3 Scenarios (1,200 events each)

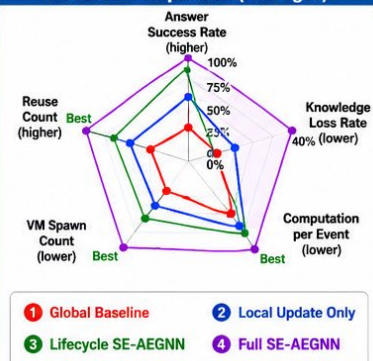
Method	Answer Success Rate (↑ higher)	Knowledge Loss Rate (↓ lower)	Computation per Event (↓ lower)	VM Spawn Count (↓ lower)	Reuse Count (↑ higher)	Active VM (final) (↓ lower)
1 Global Baseline	98.39%	25.24%	20.70	79.0	0.0	97.0
2 Local Update Only	98.31%	26.47%	4.13	77.7	0.0	95.7
3 Lifecycle SE-AEGNN	98.78%	20.87%	3.81	59.7	12.3	77.7
4 Full SE-AEGNN	100.00%	0.00%	3.14	0	30.0	18.0

Perfect Success! No Knowledge Loss! 6.59x less Computation vs Global! No New VM Spawned! Max Reuse Achieved! Minimal VM Population!

3. Time Series Examples (1,200 events)



4. Overall Comparison (Averages)



5. Ablation Study (Contribution of Each Module)

Configuration	Local Update	VM Reuse	Knowledge Recollection	Staged Delete	Knowledge Loss Rate (↓)	Computation per Event (↓)
Global Baseline	×	×	×	×	25.24%	20.70
Local Only	✓	×	×	×	26.47%	4.13
+ Reuse	✓	✓	×	×	18.56%	3.72
+ Staged Delete	✓	✓	×	✓	20.87%	3.81
Full SE-AEGNN	✓	✓	✓	✓	0.00%	3.14

All modules work synergistically for the best performance!

6. Key Takeaways

Local Update reduces computation by 6.59x	Reuse & Dormant VMs suppress unnecessary VM spawning	Knowledge Recollection eliminates knowledge loss	Staged Delete enables recovery when old topics return	Full SE-AEGNN achieves perfect success with minimal resources
---	--	--	---	---

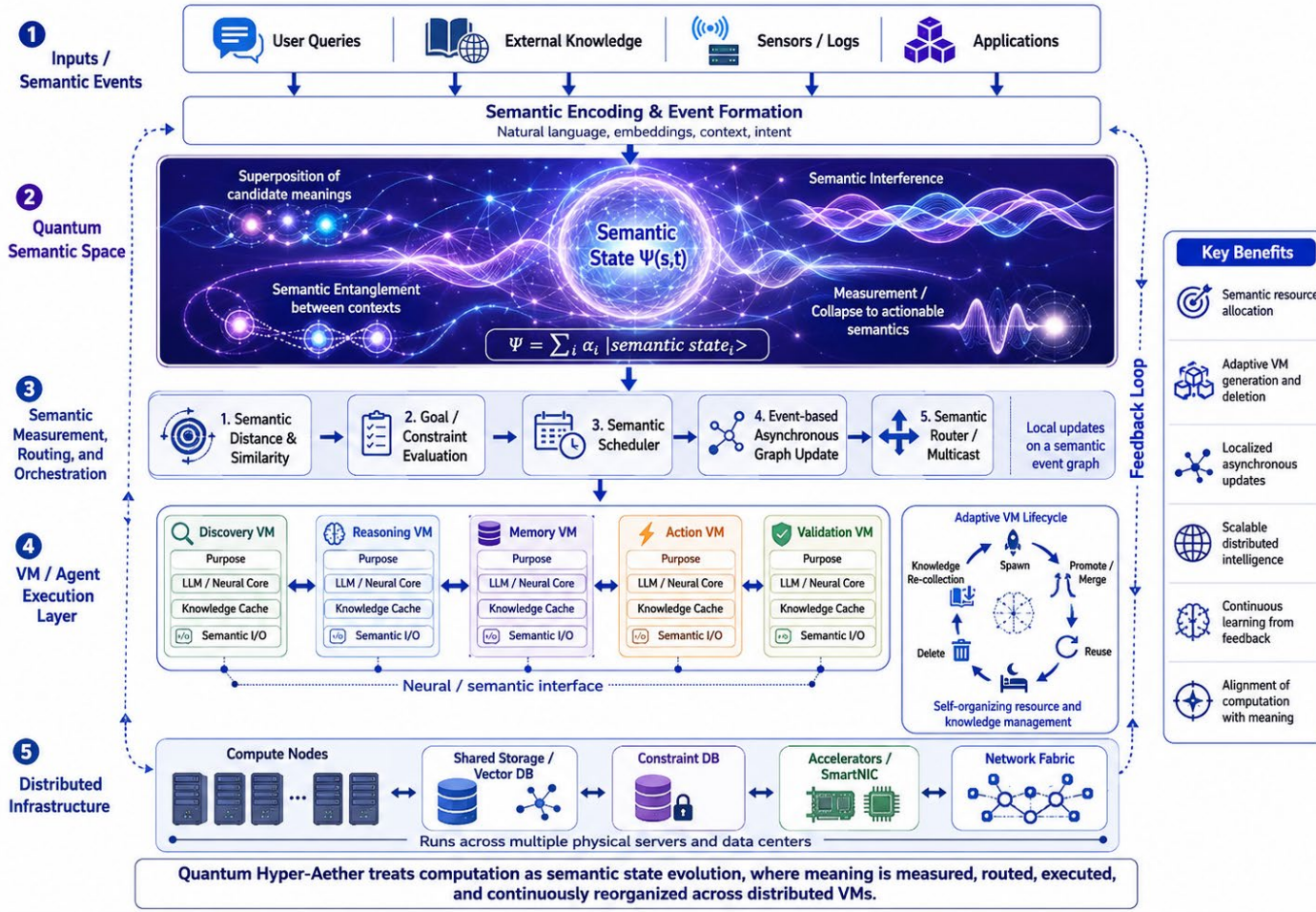
★ SE-AEGNN achieves high accuracy, zero knowledge loss, and dramatic efficiency through local update, reuse, recollection, and staged deletion.

要旨

本研究では、4つの方式を3種類の意味的嗜好シナリオに適用し、各シナリオについて1,200イベントを処理する包括的なSE-AEGNNライフサイクル・シミュレーションを実施した。比較対象は、Global Baseline、Local Update Only、Lifecycle SE-AEGNN、Full SE-AEGNNであり、評価シナリオは、嗜好が安定している場合、嗜好が徐々に変化する場合、過去の話題へ回帰する場合である。本実験では、局所グラフ更新、VM再利用、知識再収集、段階的削除が、回答品質、知識保持、計算効率、VM数の増加に与える影響を評価した。実験結果から、Full SE-AEGNNが全体として最も優れた性能を示し、回答成功率100%、知識損失率0%、1イベント当たりの最小計算量3.14を達成するとともに、最終VM数を18.0に抑えた。一方、Global Baselineでは、1イベント当たりの計算量が20.70と大きく、VM数の増加と知識損失も顕著であった。さらにアブレーション評価から、局所更新は主に計算量の削減に寄与し、VM再利用は不要なVM生成を抑制し、知識再収集は失敗または陳腐化した知識を回復し、段階的削除は過去の話題が再出現した際の知識復元を可能にすることが確認された。これらの結果は、SE-AEGNNが、局所的な非同期更新と知識を考慮したVMライフサイクル制御を統合することで、効率的・適応的・スケーラブルな意味ライフサイクル管理を実現できる可能性を示している。

Quantum Hyper-Aether: Overall Architecture

An AI-native semantic computing platform inspired by quantum concepts



Quantum Hyper-Aetherは、分散システム内における計算を、意味状態の時間的な発展として捉えるAIネイティブなセマンティック・コンピューティング・アーキテクチャです。ユーザからの質問、外部知識、センサストリーム、アプリケーションイベントなどの入力、意味表現へと符号化され、量子的概念に着想を得た意味空間へ写像されます。この空間では、複数の意味候補が同時に存在し、相互作用し、最終的に実行可能な意味状態へと選択的に収束します。

確定した意味は、意味距離評価、目標・制約評価、セマンティック・スケジューリング、イベント駆動型非同期グラフ更新、セマンティック・ルーティングから構成される階層的なオーケストレーション機構によって処理されます。実行は、Discovery VM、Reasoning VM、Memory VM、Action VM、Validation VMなどの適応型仮想マシンまたはエージェントによって行われます。各VMは、固有の目的、LLMまたはニューラル処理コア、知識キャッシュ、セマンティック入出力インターフェースを備えます。

Quantum Hyper-Aetherは、計算ノード、共有ベクトルストレージ、制約データベース、アクセラレータ、SmartNIC、ネットワークファブリックから構成される分散基盤上で動作します。適応型VMライフサイクルは、VMの生成、昇格、統合、再利用、休眠、削除、知識の再収集を支援します。VM削除によって知識が失われた場合や、ユーザの嗜好が変化した場合、システムが自律的に新しい知識を収集し、将来の問い合わせに備えることができます。

継続的なフィードバックと、意味イベントに基づく局所的な更新により、Quantum Hyper-Aetherは、拡張性、自己組織化能力、継続学習能力を備えた計算基盤を実現します。その主な目的は、計算資源、知識管理、システム動作を、エンドユーザの意味、目標、文脈の変化に適応させることです。

Quantum Hyper-Aether: 4-Method, 3-Scenario Simulation Results

Minimal integrated Python simulation comparing Global Update, Fixed VM, Lifecycle VM, and QHA + SE-AEGNN

1 Experimental Setup

- 4 methods compared
- 3 scenarios: Stationary workload, Preference drift, Knowledge deletion shock
- 30 random seeds per condition
- 600 simulation steps per run
- Metrics: response success rate, update cost per event, active VM count, retained knowledge, recovery speed

2 Main Result

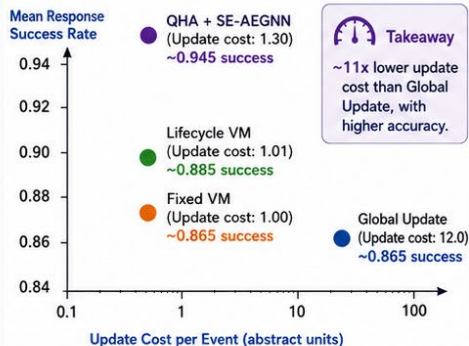


QHA + SE-AEGNN achieved the best response success rate in all three scenarios while keeping update cost **far below** Global Update.

Mean Response Success Rate

Scenario	Global Update	Fixed VM	Lifecycle VM	QHA + SE-AEGNN
Stationary workload	86.7%	86.9%	89.2%	94.7% ★
Preference drift	84.9%	85.7%	88.2%	94.6% ★
Knowledge deletion shock	86.4%	84.0%	85.7%	93.3% ★

3 Accuracy vs Update Cost



Recovery and Adaptation



Under Preference Drift

QHA + SE-AEGNN maintained the highest moving-average performance throughout the simulation.



After Knowledge Deletion Shock

Recovery steps to 90% of pre-shock level:

- Global Update: 29 steps
- Fixed VM: 64 steps
- Lifecycle VM: 78 steps
- QHA + SE-AEGNN: 34 steps



QHA + SE-AEGNN achieved near-Global-Update recovery speed with localized updates.

Additional Indicators



Mean Active VMs (approx.)

Global Update	Fixed VM	Lifecycle VM	QHA + SE-AEGNN
12.0	8.0	5.4 – 6.1	5.3 – 6.2



Mean Retained Knowledge (all scenarios)

Global Update	Fixed VM	Lifecycle VM	QHA + SE-AEGNN
0.855 – 0.875	0.840 – 0.870	0.881 – 0.903	0.909 – 0.929



Knowledge Recollection Events per Run

Global Update	Fixed VM	Lifecycle VM	QHA + SE-AEGNN
0	0	0	2.27 – 3.37



QHA + SE-AEGNN shows the highest knowledge retention and performs active re-collection.



Key Conclusion

Localized semantic-event updates, adaptive VM lifecycle control, and targeted knowledge re-collection make QHA + SE-AEGNN the **most effective and efficient** strategy in this abstract simulation model.



概要

本研究では、適応型仮想マシン、局所的な意味イベント処理、継続的な知識再編成に基づくAIネイティブなセマンティック・コンピューティング・アーキテクチャであるQuantum Hyper-Aetherの最小統合シミュレーションを提示する。比較対象として、Global Update、Fixed VM、Lifecycle VM、およびSemantic Event-based Asynchronous Graph Neural Networkを組み合わせたQuantum Hyper-Aether方式（QHA + SE-AEGNN）の4方式を評価した。実験には、定常負荷、段階的なユーザ嗜好変化、知識削除ショックの3シナリオを用いた。各条件について、30種類の乱数シードを使用し、1実行あたり600ステップのシミュレーションを行った。

実験結果から、QHA + SE-AEGNNはすべてのシナリオにおいて最も高い平均応答成功率を達成した。応答成功率は、定常負荷で94.7%、嗜好変化で94.6%、知識削除ショック後で93.3%であった。これに対し、他の方式の成功率は84.0%から89.2%の範囲であった。QHA + SE-AEGNNのイベント当たり平均更新コストは約1.30抽象単位であり、Global Updateが必要とした12.0単位を大幅に下回った。また、約5台から6台のアクティブVMで動作しながら、最も高い知識保持率を維持した。

知識削除ショックの発生後、QHA + SE-AEGNNは34シミュレーションステップでショック前性能の90%まで回復した。この回復速度はGlobal Updateの29ステップに近く、Fixed VMおよびLifecycle VMよりも大幅に高速であった。さらに、この回復はシステム全体の同期更新ではなく、対象を限定した知識再収集と局所更新によって実現された。

以上の結果は、意味イベントグラフの局所更新、適応型VMライフサイクル管理、セマンティック・ルーティング、および自律的な知識再収集を組み合わせることで、応答精度、環境変化への適応性、知識保持性能、計算効率を改善できる可能性を示している。本実験は実システムの測定ではなく、抽象的な分散イベント・シミュレーションであるが、Quantum Hyper-Aetherが、意味中心型で自己組織化し、継続的に学習する分散計算基盤となり得ることを示す初期的な証拠を提供する。



Quantum Hyper-Aether: Research Evaluation

Conceptual assessment of the architecture, simulation evidence, strengths, limitations, and next research steps

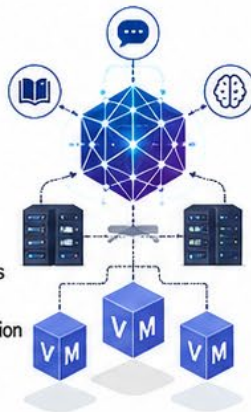


Quantum Hyper-Aether is a promising AI-native semantic computing framework with strong conceptual novelty and encouraging simulation results, but it still requires formalization and real-system validation.



1. Research Goal & Core Idea

- Meaning-centered distributed computing
- Quantum-inspired semantic state evolution
- Adaptive VM / agent lifecycle
- Localized semantic-event updates
- Autonomous knowledge re-collection



2. Main Strengths



Originality: combines semantic computing, adaptive VMs, and quantum-inspired modeling

High



System coherence: architecture spans input, semantic space, orchestration, VM execution, and infrastructure

High



Adaptivity: supports spawn, reuse, dormancy, deletion, and re-collection

High



Efficiency potential: localized updates reduce the cost of full global synchronization

High



Scalability potential: distributed VM design fits multi-node infrastructure

Promising



Learning capability: continuous feedback enables self-organization

Promising



3. Evidence from Minimal Simulation

Mean Response Success Rate (%)

Scenario	Global Update	Fixed VM	Lifecycle VM	QHA + SE-AEGNN
Stationary workload	86.7%	86.9%	89.2%	94.7%
Preference drift	84.9%	85.7%	88.2%	94.6%
Knowledge deletion shock	86.4%	84.0%	85.7%	93.3%



Update cost per event: QHA + SE-AEGNN $\approx$ 1.30 vs Global Update = 12.0



Recovery after deletion shock: 34 steps vs 29 for Global Update



Highest retained knowledge among the tested methods



4. Limitations & Open Issues

- ⚠ Current evidence is based on an abstract discrete-event simulation Needs work
- ⚠ No real deployment, hardware benchmark, or LLM latency measurement yet Needs work
- ⚠ Quantum interpretation is conceptual, not physical quantum computation Needs work
- ⚠ Mathematical formalization is incomplete for some state transitions and guarantees Needs work
- ⚠ Need validation on larger workloads, real datasets, and multi-server environments Open challenge
- ⚠ Need stronger evaluation of robustness, security, and operational complexity Open challenge



5. Overall Evaluation



Concept originality



Very High



Architectural consistency



High



Adaptivity and learning potential



High



Simulation support



Medium-High



Engineering maturity



Medium-Low



Real-world validation



Low



Overall research promise



High

Overall, Quantum Hyper-Aether is a high-potential research direction with strong conceptual contributions and encouraging early evidence, but it remains at the pre-deployment research stage.



6. Recommended Next Steps



Formalize the state-transition model



Build a larger integrated prototype



Evaluate with real semantic workloads



Measure latency, throughput, and resource cost



Test multi-node and fault-tolerant behavior



Refine knowledge deletion and re-collection policies



Key Takeaway: Strong conceptual novelty + promising simulation results + clear open research challenges.

Quantum Hyper-Aether: BC/DR Advantages

Semantic service continuity and adaptive recovery with VM-based orchestration



QHA improves BC/DR by recovering services at the level of meaning, purpose, and knowledge rather than only restoring infrastructure.

1 Why QHA Helps BC/DR

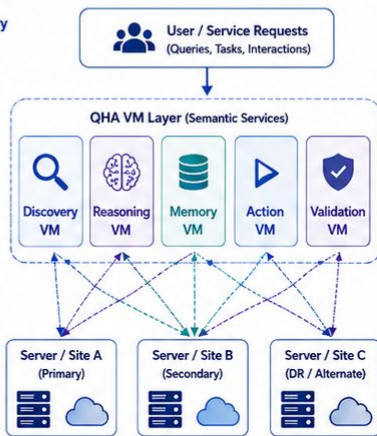
VM-based modular recovery
Recover or re-spawn only affected VMs instead of the entire system.

Semantic priority control
Restore mission-critical services first based on business meaning and user goals.

Adaptive failover
Route requests to semantically similar backup VMs when a VM fails.

Knowledge re-collection
Rebuild lost knowledge from peers, external sources, and recent interactions.

Distributed resilience
Run across multiple servers and data centers with flexible VM placement.



2 BC/DR Recovery Flow in QHA



★ Recovery is guided by service meaning, business criticality, and knowledge availability.

3 Conventional VM-DR vs QHA Semantic DR

	Conventional VM-DR	QHA Semantic DR
Recovery unit	Whole VM / application	Purpose / semantic service / VM set
Recovery priority	Static order	Context-aware semantic priority
Failover	Restore same image	Reuse, merge, re-route, or re-spawn VMs
Data recovery	Backup restoration	Backup + peer knowledge + external recollection
Objective	Return to previous state	Maintain service continuity and adapt
Operation mode	Infrastructure-centric	Meaning-centric and self-organizing

4 Key BC/DR Benefits



Lower recovery scope

Only affected semantic components need action.



Faster service continuity

Important services can resume before full restoration.



Higher flexibility

Alternate VMs can temporarily answer or process tasks.



Better knowledge resilience

Lost knowledge can be re-acquired and validated.

5 Example Semantic Recovery Policy

- 1 Safety / Compliance VM
- 2 Customer Support VM
- 3 Payment / Contract VM
- 4 Analytics VM
- 5 Exploration / Learning VM

QHA can restore services according to business meaning and continuity value.

6 Suggested BC/DR Evaluation Metrics



Semantic RTO: time until critical queries can be answered again



Knowledge RPO: amount of knowledge lost after failure



Service continuity rate: percentage of critical services available during recovery



Alternative VM success rate: success of rerouting to alternate VMs



Post-recovery answer accuracy: accuracy of responses after recovery



Re-collection cost: resources/cost to re-acquire and validate knowledge



Key Takeaway: QHA extends traditional disaster recovery into semantic service continuity by re-routing, re-spawning, and re-learning across distributed VMs.



概要

本研究では、VMベースのセマンティック・オーケストレーションを活用するQuantum Hyper-Aether (QHA) の、事業継続および災害復旧 (BC/DR) における優位性について検討する。従来のVM型災害復旧が、主としてインフラ、VMイメージ、またはアプリケーション全体の復元を目的とするのに対し、QHAは、意味、目的、知識の単位で復旧を行う。QHAでは、サービスがDiscovery VM、Reasoning VM、Memory VM、Action VM、Validation VMなどの適応型仮想マシンへ分割されるため、障害発生時には、影響を受けた意味的構成要素のみを再ルーティング、再生成、または再構成できる。

QHAは、意味的重要度に基づく優先制御、適応型フェイルオーバー、分散VM配置、知識再収集によってBC/DRを強化する。重要なサービスは、業務上の意味やユーザ目標に基づいて優先的に復旧できる。また、特定のVMが停止した場合には、意味的に近い代替VMへ処理を転送できる。失われた知識は、他のVM、外部知識源、最近の対話履歴などから部分的に再構築できるため、単純なバックアップ復元よりも柔軟な復旧が可能となる。

従来のVM-DRと比較して、QHAは、状況依存型の復旧ポリシー、動的なサービス再編成、縮退状態における継続運転を支援する。その目的は、単に障害前の状態へ戻すことではなく、意味的サービスの継続性を維持しながら、障害後の状況へ適応することである。この特性により、QHAは、知識、意図、サービスの意味がシステム運用の中心となる分散AIネイティブ・システムに適している。

QHAのBC/DR性能は、従来の指標に加え、意味的復旧時間、知識復旧時点、サービス継続率、代替VM成功率、復旧後の回答精度、知識再収集コストなどによって評価できる。以上より、QHAは、従来の災害復旧を、意味中心型かつ適応型の復旧基盤へ拡張するものであり、耐障害性を備えた分散知能システムの有望な研究方向を示している。

BC/DR Research for Generative AI Systems

Current directions, recovery targets, research gaps, and the Quantum Hyper-Aether (QHA) opportunity

1

Current Research Areas



1) Inference Continuity

- KV-cache replication
- Fast failover
- GPU / node failure recovery



2) Training Recovery

- Checkpointing
- Distributed job restart
- Pipeline reconfiguration



3) Integrity & Safety

- Model-weight integrity
- Cache corruption detection
- Semantic consistency after recovery



4) Governance & Standards

- Risk management
- Auditability
- Operational resilience

2

What Must Be Recovered?



1) Model State
weights, adapters,
quantization



2) Inference State
KV cache, active
session



3) Knowledge State
RAG documents,
embeddings, vector
index



4) Semantic State
user intent, context,
task status



5) Governance State
system prompts,
policies, guardrails



6) Tool State
external API
progress, pending
actions



7) Evaluation State
benchmarks,
quality baselines



8) Audit State
logs, provenance,
version history

3

Open Research Gaps



A) Semantic RPO / RTO

How much conversation, plan,
and user preference can be lost?



B) Service Recovery $\neq$ Semantic Recovery
Service Recovery $\neq$ Semantic Recovery



C) Consistent RAG Recovery

$$K = (D, E, I, C, V)$$

documents, embedding model, index, chunking, version

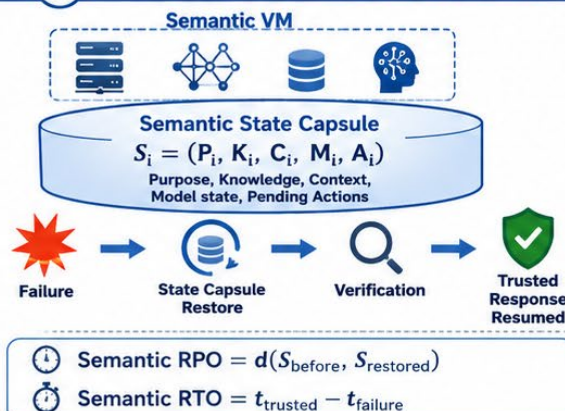


D) Autonomous Agent Recovery

Avoid duplicated or missing
actions after failure

4

QHA / Semantic-State-Aware BC/DR



Recovery Options

- 1 Restore the same VM
- 2 Generate an alternative VM
- 3 Reuse a dormant VM
- 4 Re-collect missing knowledge
- 5 Degraded operation with another LLM
- 6 Multi-VM verified recovery



Key Insight: Generative AI BC/DR must recover not only infrastructure, but also semantic state, knowledge, and trustworthy behavior.



アブストラクト

生成AIシステムにおける事業継続および災害復旧 (BC/DR) は、重要な研究領域として注目されつつある。しかし、現在の研究は、推論サービスの継続性、学習処理の復旧、モデルおよびキャッシュの完全性、ガバナンスなど、複数の領域に分散している。既存研究の多くは、GPUやノード障害からの復旧、チェックポイント、KVキャッシュの保護、運用リスク管理といった、主にインフラストラクチャ層の耐障害性を対象としている。一方、生成AIシステムのBC/DRでは、インフラストラクチャを復旧するだけでは十分ではない。モデル状態や推論状態に加えて、知識状態、意味状態、外部ツールの実行状態、評価基準、監査履歴なども一貫して復元する必要がある。この観点から、意味的RPOおよび意味的RTOの定義、サービス復旧と意味的復旧の区別、RAGパイプラインの整合性を保った復旧、自律エージェントによる処理の重複や欠落の防止などが、重要な未解決課題として挙げられる。

これらの課題に対し、本研究ではQuantum Hyper-Aether (QHA) における意味状態認識型BC/DRの可能性を示す。QHAでは、各Semantic VMが、目的、知識、文脈、モデル状態、未完了アクションから構成されるSemantic State Capsuleを保持する。障害発生時には、このカプセルを復元し、意味的一貫性を検証したうえで、信頼可能な応答処理を再開する。

この方式は、従来のBC/DRを、単なるインフラストラクチャの復旧から、意味、知識、文脈、および信頼性の継続へと拡張するものである。したがって、意味状態認識型BC/DRは、将来の高信頼な生成AIシステムを実現するための有望な基盤となる。

参考文献リストを下さい

生成AIのBC/DRに直接関係する文献と、基盤技術・標準を分けて、論文に載せやすい形式で整理します。

以下は、先ほどの図とアブストラクトに対応する**英語論文用の参考文献リスト**です。IEEEに近い形式で整理しています。

References

[1] F. Strati, S. M. Salak, S. Y. Ko, and M. Chowdhury, “DéjàVu: KV-Cache Streaming for Fast, Fault-Tolerant Generative LLM Serving,” *arXiv preprint arXiv:2403.01876*, 2024.

LLM推論時のKVキャッシュ移送と高速障害復旧を扱う代表的研究です。

[2] S. Qian et al., “Towards Resiliency in Large Language Model Serving with KevlarFlow,” *arXiv preprint arXiv:2601.22438*, 2026.

LLMサービングにおけるフェイルオーバー、KVキャッシュ保護、モデル再初期化時間の短縮を扱います。

[3] S. Jayakody et al., “GhostServe: A Lightweight Checkpointing System in the Shadow of LLM Serving,” *arXiv preprint arXiv:2605.00831*, 2026.

消失訂正符号を用いてKVキャッシュをホストメモリ上に保護し、GPU障害後の復元を高速化します。

[4] S. Nian, J. Fang, Q. Feng, Z. Wu, and F. Lai, “CacheFlow: Efficient LLM Serving with 3D-Parallel KV Cache Restoration,” *arXiv preprint arXiv:2604.25080*, 2026.

トークン、レイヤー、GPUの三方向の並列性を利用したKVキャッシュ復元方式です。

[5] W. Xu et al., “AnchorTP: Resilient LLM Inference with State-Preserving Recovery,” *arXiv preprint arXiv:2511.11617*, 2025.

モデルパラメータとKVキャッシュを保持しながらTensor Parallel構成を復旧する方式を提案しています。

[6] I. Jang, Z. Yang, Z. Zhang, X. Jin, and M. Chowdhury, “Oobleck: Resilient Distributed Training of Large Models Using Pipeline Templates,” *Proceedings of the 29th Symposium on Operating Systems Principles*, 2023.

大規模モデルの分散学習における障害耐性と高速再構成を扱う重要な研究です。

[7] S. Gandhi et al., “ReCycle: Resilient Training of Large DNNs Using Pipeline Adaptation,” *arXiv preprint arXiv:2405.14009*, 2024.

予備サーバに依存せず、障害後に学習パイプラインを再構成する方式を扱います。

[8] E. Rojas et al., “A Study of Checkpointing in Large-Scale Training of Deep Neural Networks,” *arXiv preprint arXiv:2012.00825*, 2020.

PyTorch、TensorFlowなどにおけるチェックポイント方式の性能、保存形式、スケーラビリティを評価しています。

[9] G. Wang et al., “FastPersist: Accelerating Model Checkpointing in Deep Learning,” *arXiv preprint arXiv:2406.13768*, 2024.

大規模学習におけるチェックポイント保存のI/Oボトルネックを対象としています。

[10] P. Lewis et al., “Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks,” in *Advances in Neural Information Processing Systems*, vol. 33, 2020.

RAGの基礎文献であり、モデル状態とは別に外部知識とベクトル索引を管理する考え方の根拠になります。

[11] C. Autio et al., *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*, NIST AI 600-1, National Institute of Standards and Technology, 2024.

生成AI特有のリスク、信頼性、監視、評価、インシデント管理を整理した主要な公的文書です。

[12] National Institute of Standards and Technology, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1, 2023.

AIシステムのGovern、Map、Measure、Manageというリスク管理構造を提供します。

[13] International Organization for Standardization, *ISO/IEC 42001:2023, Information Technology—Artificial Intelligence—Management System*, Geneva, Switzerland, 2023.

AIマネジメントシステムの構築、運用、監査、継続的改善に関する国際規格です。

[14] International Organization for Standardization, *ISO 22301:2019, Security and Resilience—Business Continuity Management Systems—Requirements*, Geneva, Switzerland, 2019.

一般的なBCMSの要求事項を定義する中心的な国際規格です。

[15] International Organization for Standardization, *ISO/IEC 27031:2025, Cybersecurity—Information and Communication Technology Readiness for Business Continuity*, Geneva, Switzerland, 2025.

ICT基盤を事業継続に対応させるための設計、準備、復旧の指針です。

[16] M. Swanson, P. Bowen, A. W. Phillips, D. Gallup, and D. Lynes, *Contingency Planning Guide for Federal Information Systems*, NIST Special Publication 800-34 Revision 1, National Institute of Standards and Technology, 2010.

情報システムの緊急時対応、代替設備、復旧手順、試験、保守を扱うBC/DRの基礎文献です。

QHA論文での位置付け

上記のうち、文献[1]～[9]は主に**モデル、GPU、学習ジョブ、KVキャッシュの技術的復旧**を扱います。文献[10]は**知識状態の分離と復旧**、文献[11]～[16]は**ガバナンス、BCMS、ICT継続性**の根拠になります。

一方、次の概念は、既存文献で一般的に確立された用語というより、QHA研究として提案する独自概念です。

- Semantic State Capsule
- Semantic RPO
- Semantic RTO
- Service RecoveryとSemantic Recoveryの区別
- Dormant VMを利用した意味状態復旧
- 知識再収集を含む自律的DR

したがって論文では、これらを「既存研究の引用」ではなく、**本研究の新規提案**として明確に記述するのが適切です。

QHA BC/DR Evaluation Results

4 recovery methods × 3 failure scenarios | Minimal Monte Carlo simulation | 1,000 trials per method-scenario pair (12,000 trials total)

1 Compared Recovery Methods

-  **1 Traditional Backup/Restore**
-  **2 VM Snapshot Recovery**
-  **3 Lifecycle VM Recovery**
-  **4 QHA + SE-AEGNN BC/DR**

Failure Scenarios

- Single VM Failure
- Multi-VM / Node Failure
- Knowledge Corruption + Preference Drift

 **Semantic RTO** = time until the system can again provide consistent, useful answers to critical user queries.

2 Mean Semantic RTO (minutes)

Failure Scenario	 Traditional Backup/Restore	 VM Snapshot Recovery	 Lifecycle VM Recovery	 QHA + SE-AEGNN BC/DR
Single VM Failure	25.71	11.14	7.02	3.57
Multi-VM / Node Failure	71.51	38.45	23.05	11.01
Knowledge Corruption + Preference Drift	61.66	45.97	21.19	8.97



QHA 95th-percentile Semantic RTO: 4.64 min, 14.42 min, 11.99 min

3 QHA vs Traditional Backup/Restore

 Semantic RTO reduction: 84.6–86.1%	 RPO reduction: 88.4–91.4%	 Compute cost reduction: 61.9–69.6%
 Network transfer reduction: 77.2–82.8%	 Answer success improvement: +10.4 to +32.0 points	 Semantic service continuity improvement: +39.4 to +55.5 points

 **Best overall performance: QHA + SE-AEGNN BC/DR**

4 Key Findings

-  QHA consistently achieved the lowest Semantic RTO across all failure scenarios.
-  The largest relative advantage appeared under knowledge corruption and preference drift.
-  Unlike snapshot-only recovery, QHA combines local semantic-event updates, priority-based recovery, VM regeneration, and autonomous knowledge recollection.

 **Conclusion:** QHA extends BC/DR from infrastructure restoration to semantic-service restoration.

アブストラクト

本研究では、Quantum Hyper-Aether (QHA) における事業継続および災害復旧 (BC/DR) 能力を評価するため、最小構成のモンテカルロ・シミュレーションを実施した。比較対象は、従来型バックアップ/リストア、VMスナップショット復旧、ライフサイクル型VM復旧、ならびにSemantic Event-based Asynchronous Graph Neural Network (SE-AEGNN) を統合したQHA BC/DRの4方式である。障害シナリオとして、単一VM障害、複数VMまたは物理ノード障害、知識破損とユーザー嗜好変化が同時に発生する障害の3種類を設定した。

従来のBC/DR評価では、サーバやVMが再起動した時点までを復旧完了とみなすことが多い。これに対し、本研究では、重要なユーザー質問に対して、再び整合性のある有用な回答を提供できるようになるまでの時間を「Semantic Recovery Time Objective (Semantic RTO)」として定義した。さらに、失われた意味イベント数、復旧後の回答成功率、意味的サービス継続率、復旧計算コスト、ネットワーク転送量、重要意味領域の優先復旧率、および完全復旧までのステップ数を評価指標として用いた。

各方式と各障害シナリオの組合せについて1,000回、合計12,000回の試行を行った。その結果、設定した合成パラメータの範囲では、QHA + SE-AEGNN BC/DRは従来型バックアップ/リストアと比較して、平均Semantic RTOを約84.6～86.1%短縮した。また、意味イベント損失を88.4～91.4%、復旧計算コストを61.9～69.6%、ネットワーク転送量を77.2～82.8%削減した。復旧後の回答成功率は10.4～32.0ポイント、意味的サービス継続率は39.4～55.5ポイント向上した。

特に、知識破損とユーザー嗜好変化が同時に発生するシナリオにおいて、QHAの優位性が顕著に現れた。VMスナップショット方式は過去のシステム状態を復元できる一方で、新たに必要となった知識や変化した嗜好への適応が困難である。これに対しQHAは、局所的な意味イベント更新、重要領域の優先復旧、VM再生成、および自律的な知識再収集を組み合わせることで、障害前の状態を単に再現するのではなく、現在の意味要求に適応したサービス復旧を可能にする。

以上の結果から、QHAはBC/DRの対象を、従来のインフラストラクチャ復旧から意味的サービス復旧へ拡張できる可能性を示した。ただし、本評価は合成パラメータに基づく最小シミュレーションであり、実システム性能を直接示すものではない。今後は、VM起動時間、実障害分布、ネットワーク帯域、知識再収集時間、およびLLM推論コストなどの実測値を用いた検証が必要である。

Adaptive Semantic BC/DR Simulation Results

4 recovery methods × 5 failure scenarios × 2,000 Monte Carlo runs per case (40,000 total trials)

RECOVERY METHODS (4)

-  Fixed Recovery Policy
-  Rule-Based Adaptive Recovery
-  SE-AEGNN Adaptive Recovery
-  QHA Predictive BC/DR

FAILURE SCENARIOS (5)

-  Single VM Failure
-  Physical Node Failure
-  Knowledge Corruption
-  User Preference Drift
-  Compound Cascading Failure

OVERALL PERFORMANCE SUMMARY AND RANKING

RANK	METHOD	TOTAL RECOVERY UTILITY (higher is better)	SEMANTIC RTO (s) (lower is better)	SEMANTIC RPO (lower is better)	ANSWER SUCCESS RATE (higher is better)	SEMANTIC CONTINUITY (higher is better)
1	 QHA Predictive BC/DR	0.706	30.57 s	0.157	0.725	0.792
2	 SE-AEGNN Adaptive Recovery	0.634	33.29 s	0.197	0.649	0.720
3	 Rule-Based Adaptive Recovery	0.489	49.28 s	0.321	0.503	0.554
4	 Fixed Recovery Policy	0.379	58.55 s	0.398	0.396	0.439



KEY FINDINGS

- ✓ QHA achieved the best overall recovery performance.
- ✓ QHA reduced mean Semantic RTO by about 48% versus Fixed Recovery Policy.
- ✓ QHA reduced Semantic RPO by about 61% versus Fixed Recovery Policy.
- ✓ QHA was especially strong in Knowledge Corruption, User Preference Drift, and Compound Cascading Failure.
- ✓ SE-AEGNN also outperformed rule-based and fixed recovery, showing the value of local semantic graph updates.



ADAPTIVE RECOVERY ACTIONS

-  reroute
-  snapshot restore
-  dormant VM reuse
-  knowledge recollection
-  VM regeneration
-  knowledge merge
-  hybrid recovery



INSIGHT

The results suggest that predictive semantic recovery and knowledge reuse improve both resilience and semantic service continuity. In particular, QHA combines fast recovery with lower semantic loss, making it robust under corruption, drift, and cascading failures.



Higher is better for utility, success rate, and continuity. Lower is better for RTO and RPO.

アブストラクト

本研究では、Quantum Hyper-Aether (QHA) を対象とした適応型セマンティック事業継続・災害復旧 (BC/DR) シミュレーションを提案し、固定復旧ポリシー、ルールベース適応復旧、SE-AEGNN適応復旧の3方式と比較評価した。評価には、単一VM障害、物理ノード障害、知識破損、利用者嗜好ドリフト、複合連鎖障害の5種類の障害シナリオを用いた。各方式と各シナリオの組合せについて2,000回のモンテカルロ試行を実施し、総試行数は40,000回とした。

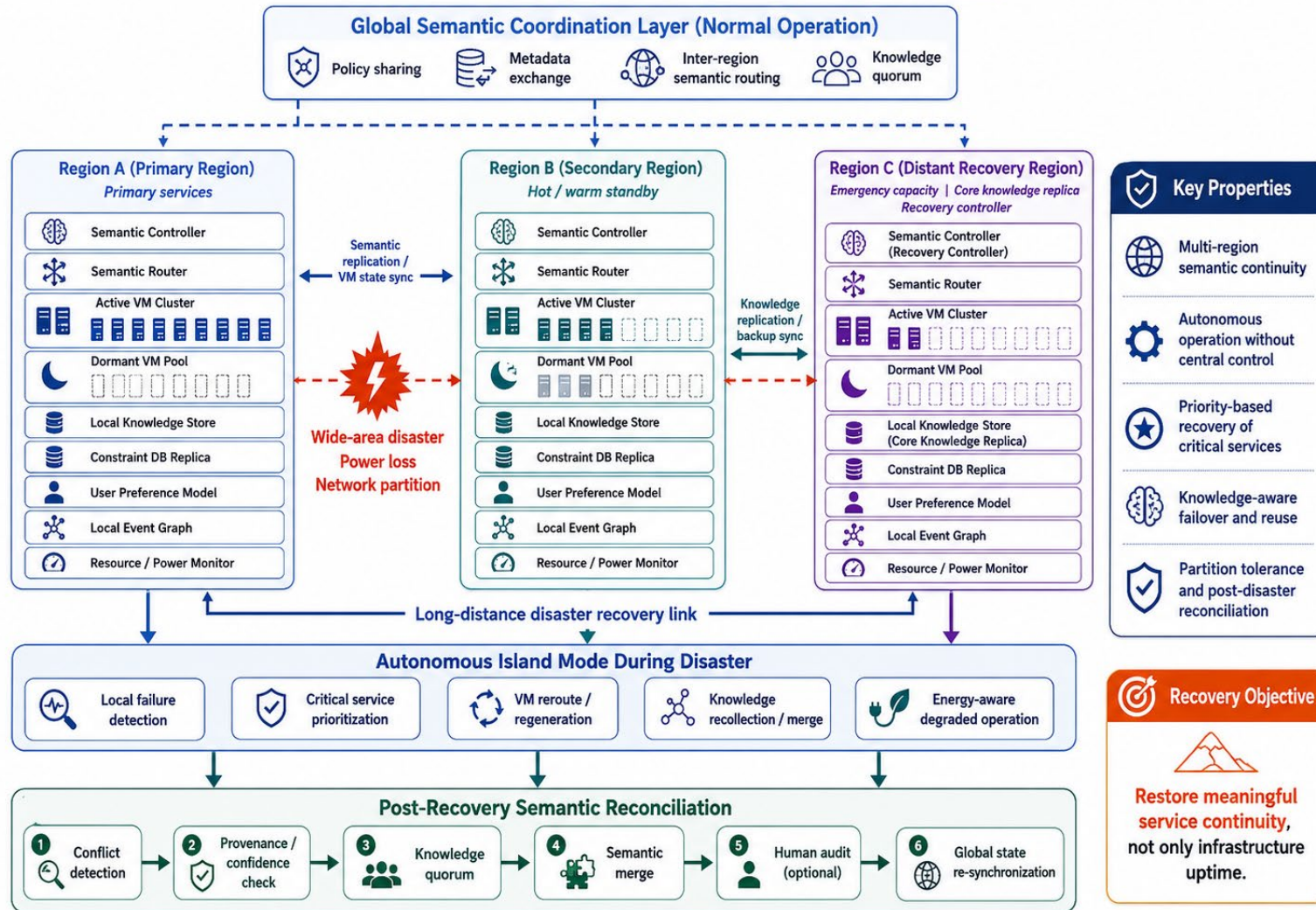
提案フレームワークでは、代替VMへのルーティング、スナップショット復旧、休眠VMの再利用、知識再収集、VM再生成、知識統合、ハイブリッド復旧を含む複数の復旧操作をモデル化した。評価指標として、Semantic Recovery Time Objective (Semantic RTO)、Semantic Recovery Point Objective (Semantic RPO)、回答成功率、意味継続率、復旧判断精度、復旧後精度、復旧振動率、総合復旧効用を用いた。

実験結果から、QHA Predictive BC/DRは全方式の中で最も高い総合性能を示した。QHAは、総合復旧効用0.706、平均Semantic RTO 30.57秒、Semantic RPO 0.157、回答成功率0.725、意味継続率0.792を達成した。固定復旧ポリシーと比較すると、平均Semantic RTOを約48%短縮し、Semantic RPOを約61%削減した。特に、知識破損、利用者嗜好ドリフト、複合連鎖障害の各シナリオにおいて、予測型セマンティック復旧と知識再利用の効果が顕著に現れた。

以上の結果は、適応的かつ予測的で、意味状態を考慮したBC/DR機構が、システムの耐障害性と意味サービス継続性を大きく改善できることを示している。本研究は仮説検証を目的としたパラメータ化シミュレーションであり、実運用環境の測定データに基づく校正が今後必要であるが、QHA型復旧戦略の有効性を示す基礎的結果を提供する。

Geo-Distributed Autonomous Semantic BC/DR: System Architecture

Autonomous multi-region semantic resilience for wide-area disasters



アブストラクト

本研究では、広域災害下で稼働するQuantum Hyper-Aether (QHA) システムを対象として、**Geo-Distributed Autonomous Semantic Business Continuity and Disaster Recovery**アーキテクチャを提案する。従来のマルチリージョンDRが主としてインフラ、仮想マシン、データベースの復旧を目的とするのに対し、本方式では、意味状態、VMの目的、知識関係、利用者嗜好モデル、サービス優先度も保存・再構成の対象とする。

提案システムは、地理的に分離されたプライマリリージョン、セカンダリリージョン、遠隔復旧リージョンから構成される。各リージョンには、セマンティックコントローラ、セマンティックルータ、稼働中および休眠状態のVMプール、ローカル知識ストア、制約情報の複製、利用者嗜好モデル、ローカルイベントグラフ、ならびに資源・電力監視機構を配置する。通常時には、グローバルセマンティック協調層が、ポリシー共有、メタデータ交換、リージョン間ルーティング、状態同期、知識クォーラム形成を支援する。広域災害によって停電、リージョン障害、通信分断が発生した場合、各生存リージョンは自律アイランドモードへ移行する。このモードでは、中央制御に依存せず、ローカルコントローラが障害検出、重要サービスの優先順位付け、VMの再ルーティングまたは再生成、休眠資源の再利用、知識の再収集・統合、電力制約を考慮した縮退運転を自律的に実行する。これにより、システム全体の持続性が失われた状況でも、重要な意味サービスを継続できる。通信復旧後は、競合検出、知識の出所および信頼度の評価、知識クォーラム形成、セマンティックマージ、必要に応じた人間監査、グローバル状態の再同期を通じて、リージョン間の意味状態を再統合する。したがって、本方式の復旧目標は、単なるインフラ可用性の回復ではなく、意味的かつ文脈的に適切なサービス継続性の復元にある。本提案は、深刻な分散障害下でも、意味的一貫性、リージョン単位の自律運転、適応的なサービス復旧を維持できるAIネイティブクラウド基盤の実現に向けた基礎を提供する。

Geo-Distributed Autonomous Semantic BC/DR Simulation Results

3-region model × 4 recovery methods × 5 wide-area disaster scenarios × 250 Monte Carlo runs per case (5,000 total runs)

RECOVERY METHODS (4)

- 1 Conventional Multi-Region DR
- 2 Active-Standby Geo-DR
- 3 Autonomous Island BC/DR
- 4 QHA Geo-Distributed Semantic BC/DR

WIDE-AREA DISASTER SCENARIOS (5)

- 1 Primary Region Loss
- 2 Dual-Region Network Partition
- 3 Regional Power Shortage
- 4 Knowledge Replica Corruption
- 5 Wide-Area Cascading Disaster

OVERALL PERFORMANCE SUMMARY

Rank	Method	Geo Recovery Utility (higher is better)	Global Semantic RTO (lower is better)	Critical Service Survival Rate (higher is better)	Semantic Partition Tolerance (higher is better)	Knowledge Divergence (lower is better)	Reconciliation Accuracy (higher is better)
1	QHA Geo-Distributed Semantic BC/DR	0.615	40.66	0.695	0.582	0.069	0.889
2	Autonomous Island BC/DR	0.527	50.45	0.649	0.550	0.117	0.702
3	Active-Standby Geo-DR	0.402	65.00	0.561	0.494	0.187	0.495
4	Conventional Multi-Region DR	0.363	65.00	0.528	0.474	0.225	0.403

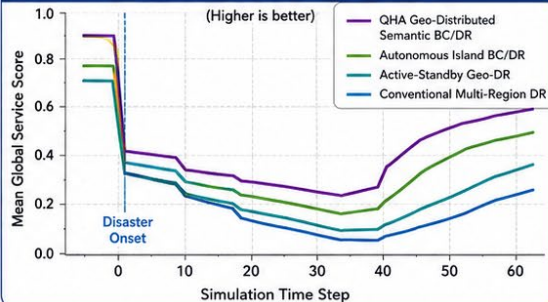
KEY FINDINGS

- ✓ QHA achieved the best overall geo-distributed recovery performance.
- ✓ QHA showed the lowest knowledge divergence and the highest reconciliation accuracy.
- ✓ Autonomous Island BC/DR outperformed both conventional and active-standby DR.
- ✓ The results support the value of autonomous regional operation, dormant VM reuse, and semantic reconciliation.
- ✓ QHA remained strongest under severe wide-area cascading failures.

QHA VS CONVENTIONAL DR

- +69%** Geo Recovery Utility
- ~37%** lower Global Semantic RTO
- ~69%** lower Knowledge Divergence
- Reconciliation Accuracy: 0.889 vs 0.403**

WIDE-AREA CASCADING DISASTER



INSIGHT

The simulation suggests that geo-distributed semantic recovery is more effective when each region can operate autonomously during partitions and later reconcile semantic knowledge safely after reconnection. QHA combines these capabilities most successfully.

アブストラクト

本研究では、広域災害下におけるQuantum Hyper-Aetherシステムを対象として、Geo-Distributed Autonomous Semantic Business Continuity and Disaster Recoveryアーキテクチャを評価する。提案方式は、従来のマルチリージョンDRを拡張し、インフラ可用性だけでなく、意味状態、知識の健全性、サービス目的、利用者嗜好、通信分断後の再統合も復旧対象として扱う。シミュレーションでは、プライマリリージョン、セカンダリリージョン、遠隔復旧リージョンから成る3リージョンモデルを構築した。比較対象は、Conventional Multi-Region DR、Active-Standby Geo-DR、Autonomous Island BC/DR、QHA Geo-Distributed Semantic BC/DRの4方式である。評価シナリオとして、Primary Region Loss、Dual-Region Network Partition、Regional Power Shortage、Knowledge Replica Corruption、Wide-Area Cascading Disasterの5種類を設定した。各方式と各シナリオの組合せについて250回のモンテカルロ試行を実施し、総試行数は5,000回とした。モデルには、リージョンごとのサービス容量、休眠VMの起動、電力低下、ネットワーク分断、知識損傷、自律アイランド運転、知識再収集、ならびに復旧後のセマンティック再統合を含めた。評価指標として、Global Semantic Recovery Time Objective、Semantic Recovery Point Objective、Critical Service Survival Rate、Semantic Partition Tolerance、Knowledge Divergence、Reconciliation Accuracy、Energy-Constrained Continuity、Semantic Continuity、Geo Recovery Utilityを用いた。

実験結果から、QHA Geo-Distributed Semantic BC/DRは0.615のGeo Recovery Utilityを達成し、Autonomous Island BC/DRの0.527、Active-Standby Geo-DRの0.402、Conventional Multi-Region DRの0.363を上回った。また、QHAはKnowledge Divergenceを0.069まで抑え、Reconciliation Accuracy 0.889を達成した。平均Global Semantic RTOは40.66シミュレーションステップであり、従来方式およびActive-Standby方式の65.00ステップより短かった。Autonomous Island BC/DRも従来2方式を上回り、通信分断時における分散型リージョン自律運転の重要性が示された。

以上の結果は、各リージョンが重要な意味サービスを独立して継続し、通信回復後に分散知識を安全に再統合できる場合、広域災害に対する耐障害性が向上することを示している。特に、自律的リージョン制御、休眠VM再利用、知識を考慮した復旧、セマンティック再統合の組合せにより、QHAは深刻な連鎖障害に対して高い堅牢性を示した。なお、本研究の数値パラメータは実運用測定値ではなく概念モデルに基づくため、今後は実際のリージョン障害データを用いた校正が必要である。

Why Semantic-Unit Processing Fits Geo-Distributed Systems

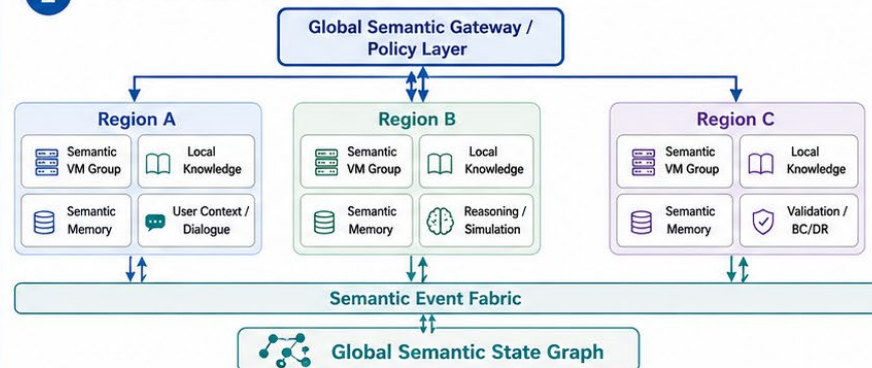
Result Summary for Geo-Distributed Autonomous Semantic Processing

1 Core Idea



- Process, route, and recover by semantic units rather than only packets or VMs
- Move meaning and intent, not always the full raw data
- Use semantic state, context, purpose, and knowledge as control targets
- Coordinate distributed execution through a global semantic state graph

2 Geo-Distributed Semantic Fabric



3 Why Better Than Conventional WAN-Centric Control

	Conventional	Semantic-Unit Approach
Unit	Packet / Request	Meaning / Task / Intent
Routing	Nearest / least loaded	Best semantic match + availability
Data movement	Move raw data	Move summaries / embeddings / context
Recovery	Bit-level replica restore	Meaning-level functional recovery
Coordination	Tight synchronization	Asynchronous semantic events

4 Key Advantages in Wide-Area Distributed Processing



Conclusion: Semantic-unit processing enables efficient routing, resilient recovery, and adaptive coordination across geo-distributed regions.

Conceptual view aligned with Quantum Hyper-Aether / semantic VM architecture

概要

セマンティック単位の処理は、パケット、リクエスト、仮想マシンだけに依存するのではなく、意味、意図、文脈、知識を主要な制御対象として扱うため、広域分散システムに適している。本方式では、分散実行をセマンティック状態とグローバルなセマンティック状態グラフによって協調させる。これにより、常に生データ全体を転送するのではなく、要約、埋め込み表現、意図などを移動させることができる。その結果、広域ネットワークの通信量を削減し、ルーティング品質を向上させるとともに、知識や機能に基づくリージョン単位の専門化を実現できる。従来のWAN中心型制御と比較して、セマンティック単位の処理は、知識を考慮したルーティング、非同期協調、意味レベルの復旧を可能にする。これにより、システムは障害に対してより強くなり、適応型BC/DRとの整合性も高まる。したがって、広域分散セマンティック処理は、拡張可能な協調、セマンティック整合性、複数リージョンにまたがる段階的な性能劣化への対応を必要とする大規模自律システムに対して、効率的かつ堅牢な基盤を提供する。