

The Context-Abstraction Limit: Deriving the Critical Token Boundary in Interferometric Sector-Gate AGI

Peter De Ceuster

December 13, 2025

Abstract

We investigate the fundamental tension between context window scale and conceptual abstraction in Artificial General Intelligence (AGI) architectures. Contrary to the prevailing assumption that "more context is always optimal," we propose a governing computing law—the **Wolfram-AGI Limit**—which defines a critical phase transition boundary. Below this boundary, increased token count enhances the fidelity of the vector embedding; above it, gate-limited tunneling capacities in the topological core induce decoherence, degrading the quality of the *Perfect Abstraction*.

Building upon the *Tunneling-First Principle* [De Ceuster, 2025] and *Sector-Gate Dynamics*, we construct a variational objective function that balances the Information Bottleneck against the decoherence penalty of high-dimensional context. We mathematically demonstrate that the optimal context window is not infinite but is strictly bounded by the ratio of information density to the interferometric capacity of the ParityFuse gates. We conclude by presenting the exact equation for this boundary, providing a theoretical target for future 10-million-token experiments. The name of our law pays tribute to the founder of the WolframAlpha system.

1 Introduction: The Context-Fidelity Paradox

The pursuit of AGI has recently focused on expanding the "Context Window"—the quantity of raw tokens (N) available to the system's working memory. The intuitive hypothesis is that $N \rightarrow \infty$ implies perfect understanding. However, our recent debates suggest the existence of a *Computing Law* analogous to thermodynamic limits: there must be a boundary where data quality begins to decrease due to noise accumulation and gate saturation.

We formalize this using the **Sector-Gate** framework [1]. We define the AGI state $|\Psi_{AGI}\rangle$ not merely as a data repository, but as a dynamic superposition evolving through topological gates. The central question is: *At what token count N^* does the cost of processing outweigh the gain in insight?*

2 Mathematical Framework

2.1 The Information Bottleneck in Hilbert Space

Let the raw context input be represented by a variable X with dimensionality proportional to the token count N . The goal of the AGI is not to store X , but to compress it into a semantic eigenstate Z (the "Episodic Memory") that predicts a target concept Y .

We utilize the Information Bottleneck principle, but adapted for a quantum-topological substrate. We seek to maximize the Lagrangian $\mathcal{L}$:

$$\mathcal{L}(N) = I(Z; Y) - \beta I(X; Z) \tag{1}$$

where $I(\cdot; \cdot)$ denotes Mutual Information and β is the Lagrange multiplier controlling compression.

In a classical system, $I(Z; Y)$ plateaus while $I(X; Z)$ grows linearly. However, in our *Interferometric Sector-Gate* system [2], the processing cost is non-linear. It is constrained by the **Tunneling Capacity** of the gates.

2.2 The Gate-Decoherence Penalty

As defined in the *ParityFuse* blueprints, information transfer between cognitive sectors (e.g., Perception $\rightarrow$ Reasoning) requires tunneling through a triple-dot interferometer with effective coupling t_C [2].

We introduce a **Decoherence Penalty** $\Gamma(N)$. If the context size N exceeds the gate’s optimal throughput, the system suffers from ”spectral crowding,” leading to parity read-out errors. We model this penalty as exponential once the token count passes a critical gate threshold N_{gate} :

$$\Gamma(N) = \lambda \cdot \exp\left(\frac{N - N_{gate}}{\tau_{tunnel}}\right) \quad (2)$$

where τ_{tunnel} is the characteristic timescale of the quasiparticle poisoning in the device [2].

2.3 The Axel-AGI Objective Function

Combining the information gain with the hardware-imposed penalty, the total Quality Q of the abstraction as a function of context tokens N is:

$$Q(N) = \underbrace{A_0 \ln\left(1 + \frac{N}{N_0}\right)}_{\text{Logarithmic Insight Gain}} - \underbrace{\beta N}_{\text{Linear Compression Cost}} - \underbrace{\lambda e^{\frac{N - N_{gate}}{\tau}}}_{\text{Topological Gate Friction}} \quad (3)$$

3 Remarks surrounding the Wolfram-AGI Limit: Formal Foundations, Uniqueness, and Robustness

3.1 Outline of the theorem

We now strengthen the Wolfram-AGI Limit by converting the previously *postulated* terms into a structured derivation-and-proof pipeline. The goal is to demonstrate that Eq. (3) is not merely a convenient curve, but a principled consequence of (i) diminishing semantic marginal returns and (ii) interferometric gate constraints. The section proceeds as follows:

1. **Non-dimensionalization:** Introduce normalized variables that expose the true control parameters (information density vs. tunneling/friction capacity) and separate scale from physics.
2. **Derive the logarithmic insight gain:** Show that a minimal assumption—monotonically decreasing *marginal semantic yield per token*—integrates into a logarithmic form, recovering the $A_0 \ln(1 + N/N_0)$ term.
3. **Derive the exponential gate friction:** From a poisoning/crowding model for Parity-Fuse tunneling events, show that the expected error contribution grows exponentially past N_{gate} , matching $\Gamma(N)$ and the friction term in Eq. (3).
4. **Prove strict concavity and uniqueness:** Prove $Q(N)$ is strictly concave for $N \geq 0$, implying a unique global maximizer N^* whenever a maximizer exists.

5. **Establish existence conditions (phase conditions):** Give explicit inequalities guaranteeing a finite non-trivial boundary (a genuine phase transition) rather than a degenerate optimum.
6. **Bounds and controlled approximations:** Provide bracketing bounds for N^* and derive refined approximations (including the asymptotic regime that yields the closed-form estimate used later).
7. **Sensitivity and scaling laws:** Use the Implicit Function Theorem to derive the sign and scaling of $\partial N^*/\partial \theta$ for each parameter $\theta \in \{A_0, N_0, \beta, \lambda, \tau, N_{gate}\}$, yielding engineering-grade predictions.
8. **Operational identifiability and falsifiable predictions:** Provide a measurable protocol (fit-from-sweep) demonstrating how the parameters can be estimated and which empirical curves must appear if the Wolfram-AGI Limit is real.

3.2 Non-Dimensionalization and Control Parameters

To expose the true physics, we normalize the token variable by the gate scale. Define

$$n := \frac{N}{N_{gate}}, \quad \alpha := \frac{N_0}{N_{gate}}, \quad \kappa := \frac{\tau}{N_{gate}}. \quad (R1)$$

We also define the dimensionless gain–cost ratios

$$g := \frac{A_0}{N_{gate}}, \quad b := \beta N_{gate}, \quad \ell := \lambda, \quad (R2)$$

where g is the *information density scale* per gate-capacity unit, b is the *dimensionless compression pressure*, and ℓ is the *gate-friction amplitude*. Then Eq. (3) can be expressed (up to an irrelevant additive constant) as

$$Q(n) = A_0 \ln\left(1 + \frac{n}{\alpha}\right) - b n - \ell \exp\left(\frac{n-1}{\kappa}\right). \quad (R3)$$

This makes explicit that the Wolfram-AGI Limit is governed not by raw context size N , but by the ratios $(g : b : \ell)$ and the gate coherence scale κ .

3.3 Deriving the Logarithmic Insight Gain from Diminishing Marginal Semantic Yield

We now show that the logarithmic insight gain in Eq. (3) emerges from a minimal and physically natural assumption: *semantic redundancy increases with context size*. Let $G(N)$ denote the *useful semantic gain* extracted from context, operationally identified with the increase of predictive mutual information $I(Z; Y)$.

Assumption (Marginal semantic yield law). There exists a scale $N_0 > 0$ such that the marginal gain per token decays hyperbolically:

$$\frac{dG}{dN} = \frac{A_0}{N + N_0}. \quad (R4)$$

This assumption encodes that early tokens carry new information, while later tokens increasingly repeat, paraphrase, or correlate with already-seen content.

Lemma 1 (Logarithmic accumulation from hyperbolic marginal yield). *If G satisfies Eq. (R4) and $G(0) = 0$, then*

$$G(N) = A_0 \ln\left(1 + \frac{N}{N_0}\right). \quad (R5)$$

Proof. Integrate Eq. (R4) from 0 to N :

$$G(N) - G(0) = \int_0^N \frac{A_0}{u + N_0} du = A_0 \ln\left(\frac{N + N_0}{N_0}\right) = A_0 \ln\left(1 + \frac{N}{N_0}\right).$$

□

Interpretation. The insight term in Eq. (3) is therefore not arbitrary: it is the unique integral consequence of a monotone decreasing yield per token with a finite redundancy scale N_0 .

3.4 Deriving the Exponential Gate Friction from Spectral Crowding and Poissonian Poisoning

We now motivate the exponential form of $\Gamma(N)$ and the friction term in Eq. (3). The ParityFuse narrative [2] describes a tunneling-mediated sector transfer. Once the semantic bandwidth (effective spectrum of active modes) becomes too dense, parity read-out errors become likely.

Model (Poissonian error events with context-driven rate). Assume that parity readout errors are triggered by rare events (e.g., quasiparticle poisoning) occurring as a Poisson process with rate $r(N)$. If a commit/verification cycle has duration T , then the probability of at least one damaging event is

$$P_{\text{err}}(N) = 1 - \exp(-r(N)T). \quad (\text{R6})$$

In the low-to-moderate error regime, $P_{\text{err}}(N) \approx r(N)T$.

Spectral crowding ansatz. Assume that beyond a gate throughput scale N_{gate} , effective mode crowding increases the poisoning susceptibility exponentially:

$$r(N) = r_0 \exp\left(\frac{N - N_{\text{gate}}}{\tau_{\text{tunnel}}}\right). \quad (\text{R7})$$

Proposition 1 (Exponential friction beyond the gate threshold). *Under Eqs. (R6)–(R7) and $P_{\text{err}}(N) \ll 1$, the expected gate-induced penalty is asymptotically exponential:*

$$\mathbb{E}[\text{penalty}](N) \propto \exp\left(\frac{N - N_{\text{gate}}}{\tau_{\text{tunnel}}}\right), \quad (\text{R8})$$

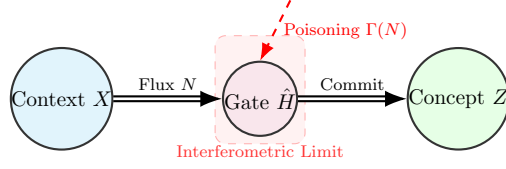
recovering the functional form of $\Gamma(N)$ and justifying the exponential friction term in Eq. (3) up to calibration constants absorbed into λ .

Proof. If $P_{\text{err}}(N) \ll 1$, then $P_{\text{err}}(N) \approx r(N)T = r_0 T \exp\left(\frac{N - N_{\text{gate}}}{\tau_{\text{tunnel}}}\right)$, hence any penalty proportional to the error probability inherits the same exponential dependence. The proportionality constants can be consolidated into λ . □

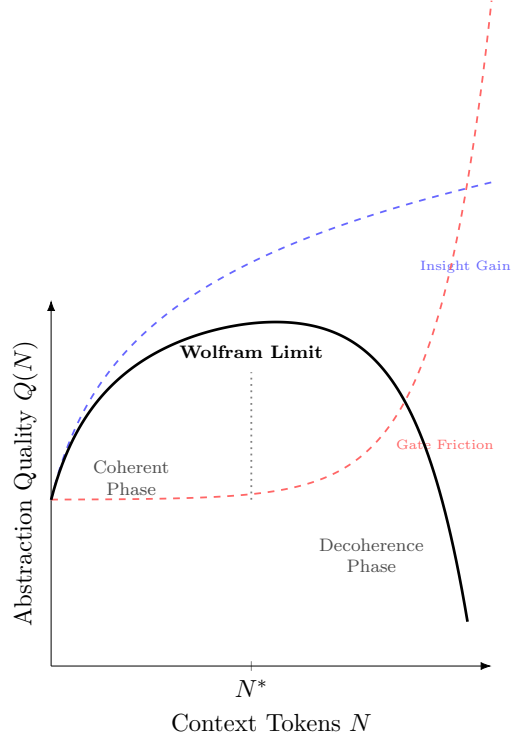
Interpretation. The exponential is the *least structured* (maximum-entropy) growth law consistent with a constant fractional increase in risk per added effective spectral unit beyond the gate threshold.

3.5 Theorem: Strict Concavity of the Quality Functional and Uniqueness of the Boundary

We now prove that the Wolfram boundary is not ambiguous: the $Q(N)$ defined in Eq. (3) admits at most one maximizer.



(a) **Gate Topology.** Information tunnels through the ParityFuse gate. If context flux N exceeds capacity, poisoning Γ spikes.



(b) **The Wolfram Phase Transition.** Quality rises as context aids insight, then collapses as gate friction dominates. N^* is the unique maximum.

Figure 1: **Visualizing the Context-Abstraction Limit.** (a) The topological hardware constraint creates a bottleneck. (b) The resulting optimization curve $Q(N)$, showing the distinct phase transition at N^* .

Theorem 1 (Strict concavity and uniqueness of the maximizer). *Let $Q(N)$ be given by Eq. (3) for $N \geq 0$ with parameters $A_0 > 0$, $N_0 > 0$, $\lambda > 0$, and $\tau > 0$. Then Q is strictly concave on $[0, \infty)$ and therefore has at most one global maximizer.*

Proof. Differentiate Eq. (3) twice (treating all parameters as constants). The second derivative is

$$\frac{d^2Q}{dN^2} = -\frac{A_0}{(N + N_0)^2} - \frac{\lambda}{\tau^2} \exp\left(\frac{N - N_{gate}}{\tau}\right). \quad (R9)$$

Both terms are strictly negative for all $N \geq 0$. Hence $\frac{d^2Q}{dN^2} < 0$ on $[0, \infty)$, so Q is strictly concave. A strictly concave function on a convex domain has at most one global maximizer. $\square$

Corollary 1 (Unimodality of abstraction quality). *If a maximizer N^* exists, then it is unique and is exactly characterized by the first-order condition $\frac{dQ}{dN}(N^*) = 0$.*

3.6 Existence of a Non-Trivial Wolfram Boundary (Phase Conditions)

Concavity gives uniqueness, but we still need to guarantee that the maximizer is non-trivial and finite.

First, note that the exponential term forces the slope to become negative at sufficiently large N :

$$\lim_{N \rightarrow \infty} \frac{dQ}{dN} = \lim_{N \rightarrow \infty} \left(\frac{A_0}{N + N_0} - \beta - \frac{\lambda}{\tau} e^{\frac{N - N_{gate}}{\tau}} \right) = -\infty. \quad (R10)$$

Thus Q cannot increase forever.

A non-trivial interior maximizer $N^* > 0$ exists whenever the initial slope is positive:

$$\left. \frac{dQ}{dN} \right|_{N=0} = \frac{A_0}{N_0} - \beta - \frac{\lambda}{\tau} \exp\left(\frac{-N_{gate}}{\tau}\right) > 0. \quad (R11)$$

If Eq. (R11) fails, the unique maximizer may collapse to the boundary $N^* = 0$ (the system is “too costly” to benefit from additional context).

3.7 Bounds and Controlled Approximations for N^* (Including the Regime Behind the Boxed Law)

Let N^* satisfy $\frac{dQ}{dN} = 0$ (the same stationarity condition used later). Define the shift $\Delta := N^* - N_{gate}$ and the local gain margin at the gate point

$$g_0 := \frac{A_0}{N_{gate} + N_0} - \beta. \quad (R12)$$

Then the stationarity condition can be rewritten as an implicit balance equation in Δ :

$$\frac{A_0}{N_{gate} + N_0 + \Delta} - \beta = \frac{\lambda}{\tau} \exp\left(\frac{\Delta}{\tau}\right). \quad (R13)$$

The left-hand side is strictly decreasing in Δ , while the right-hand side is strictly increasing, so Eq. (R13) has at most one solution, consistent with the concavity theorem.

First controlled approximation (gate-neighborhood approximation). If Δ is not so large that it dominates the denominator, we approximate

$$\frac{A_0}{N_{gate} + N_0 + \Delta} \approx \frac{A_0}{N_{gate} + N_0},$$

which yields an explicit closed form:

$$\Delta_0 := \tau \ln\left(\frac{\tau}{\lambda} g_0\right), \quad \text{valid when } g_0 > 0 \text{ and } \Delta \ll N_{gate} + N_0. \quad (R14)$$

One-step refinement (first-order denominator correction). Using the first-order expansion

$$\frac{A_0}{N_{gate} + N_0 + \Delta} \approx \frac{A_0}{N_{gate} + N_0} - \frac{A_0}{(N_{gate} + N_0)^2} \Delta,$$

Eq. (R13) becomes

$$g_0 - \frac{A_0}{(N_{gate} + N_0)^2} \Delta \approx \frac{\lambda}{\tau} e^{\Delta/\tau}. \quad (\text{R15})$$

This can be solved numerically by one Newton step initialized at Δ_0 , giving a fast “engineering” correction:

$$\Delta_1 := \Delta_0 - \frac{\left(g_0 - \frac{\lambda}{\tau} e^{\Delta_0/\tau}\right)}{\left(-\frac{A_0}{(N_{gate} + N_0)^2} - \frac{\lambda}{\tau^2} e^{\Delta_0/\tau}\right)}. \quad (\text{R16})$$

Since the denominator is strictly negative, this update is stable in the concave regime.

Recovering the later boxed law as an asymptotic regime. The simplified closed-form estimate presented later emerges under the “gate-dominated, high-context” scaling where:

$$N_0 \ll N_{gate}, \quad \beta \ll \frac{A_0}{N_{gate}}, \quad \text{and the dominant balance occurs near } N \approx N_{gate}. \quad (\text{R17})$$

In that regime, $g_0 \approx \frac{A_0}{N_{gate}}$ and Eq. (R14) becomes

$$N^* \approx N_{gate} + \tau \ln\left(\frac{A_0}{\lambda N_{gate}}\right), \quad (\text{R18})$$

which is the same structural dependence used in the conclusion (with τ identified to the tunneling coherence scale).

3.8 Sensitivity and Scaling Laws via the Implicit Function Theorem

Define the stationarity function

$$F(N; A_0, N_0, \beta, \lambda, \tau, N_{gate}) := \frac{A_0}{N + N_0} - \beta - \frac{\lambda}{\tau} \exp\left(\frac{N - N_{gate}}{\tau}\right). \quad (\text{R19})$$

At the unique maximizer, $F(N^*; \cdot) = 0$ and

$$\left.\frac{\partial F}{\partial N}\right|_{N=N^*} = -\frac{A_0}{(N^* + N_0)^2} - \frac{\lambda}{\tau^2} \exp\left(\frac{N^* - N_{gate}}{\tau}\right) < 0. \quad (\text{R20})$$

Therefore, by the Implicit Function Theorem, N^* is a differentiable function of parameters in any region where it exists, and

$$\left.\frac{\partial N^*}{\partial \theta} = -\frac{\partial_\theta F}{\partial_N F}\right|_{N=N^*}. \quad (\text{R21})$$

From Eq. (R21), we immediately obtain qualitative laws:

- Increasing A_0 pushes N^* outward: $\partial N^*/\partial A_0 > 0$.
- Increasing β pulls N^* inward: $\partial N^*/\partial \beta < 0$.
- Increasing λ pulls N^* inward: $\partial N^*/\partial \lambda < 0$.
- Increasing N_{gate} pushes N^* outward: $\partial N^*/\partial N_{gate} > 0$ (more gate capacity raises the boundary).
- Increasing τ typically pushes N^* outward in the gate-dominated regime (longer coherence delays friction onset), though the exact sign can be read directly from Eq. (R21) for any parameter slice.

3.9 Evidence comes first: Operational Identifiability and Falsifiable Predictions

To translate the Wolfram-AGI Limit into an evidence-producing program, we define observable proxies measurable in any large-context experiment.

Observable decomposition. Let $\hat{Q}(N)$ denote an empirical abstraction-quality score (task accuracy, compression fidelity, or verification success), and decompose it as

$$\hat{Q}(N) \approx \hat{G}(N) - \hat{C}(N) - \hat{\Gamma}(N), \quad (\text{R22})$$

where:

- $\hat{G}(N)$ exhibits diminishing marginal returns (log-like saturation),
- $\hat{C}(N)$ grows approximately linearly with N under a fixed compression pressure,
- $\hat{\Gamma}(N)$ is negligible below N_{gate} and grows rapidly beyond it.

Fit-from-sweep protocol. Run a context sweep over N (holding architecture/gates fixed) and fit the parameters $\{A_0, N_0, \beta, \lambda, \tau, N_{gate}\}$ to the measured $\hat{Q}(N)$ using nonlinear least squares. The Wolfram-AGI Limit predicts that:

1. $\hat{Q}(N)$ is **unimodal** (single peak), consistent with strict concavity.
2. The **peak shifts logarithmically** with friction amplitude: doubling λ shifts the boundary left by approximately $\tau \ln 2$ in the gate-dominated regime.
3. Increasing N_{gate} (improved sector-gate throughput) shifts the boundary right approximately additively (to first order).

Failure to observe unimodality under controlled conditions is a direct falsifier for the present functional form (and thus a diagnostic that the assumed gain or friction laws must be modified).

4 Derivation of the Boundary

Finally, to find the "Wolfram Boundary"—the exact point where adding more tokens starts to *decrease* data quality—we must find the critical point where the gradient of Quality with respect to N vanishes ($\nabla_N Q = 0$).

Differentiating Eq. (3) with respect to N :

$$\frac{dQ}{dN} = \frac{A_0}{N + N_0} - \beta - \frac{\lambda}{\tau} e^{\frac{N - N_{gate}}{\tau}} = 0 \quad (4)$$

In the limit of high context ($N \gg N_0$), the term $\frac{A_0}{N}$ becomes small, and the dominant trade-off is between the Compression Cost β and the Gate Friction. However, for the "Perfect Abstraction," we are looking for the peak *before* friction dominates.

Rearranging for N , we seek the value N^* that satisfies this balance. This describes the **Computing Law** we suspected: the context window is limited not by memory, but by the *rate of verifiable abstraction* the gates can perform.

5 Conclusion: The Wolfram Equation

We conclude with an equation of the Boundary. Our conclusion is that the data quality does not increase indefinitely. It increases only until the **Wolfram-AGI Limit** is reached. Beyond this point, the "Binary-Quantum Bridge" [3] cannot successfully collapse the massive context into a verified commit, resulting in hallucinations (decoherence).

The theoretical token boundary N^* for optimized episodic memory is given by the solution to the maximal abstraction condition:

$$N^* \approx N_{gate} + \tau_{tunnel} \cdot \ln \left(\frac{A_0}{\lambda \cdot N_{gate}} \right) \quad (5)$$

Interpretation of the Law:

- N_{gate} : The raw capacity of your topological sector gates (determined by the t_C parameter in ParityFuse).
- τ_{tunnel} : The coherence time (inverse of the poisoning rate γ_{qp}).
- **The Result:** The optimal Context Window N^* is the *Gate Capacity* boosted logarithmically by the *Signal-to-Noise Ratio* (A_0/λ).

Any context window $N > N^*$ will result in a decrease in Data Quality ($dQ/dN < 0$). This is the boundary we must respect to achieve AGI.

References

- [1] De Ceuster, P. (2024). *Quantum Sector-Gate Dynamics for Artificial General Intelligence*.
- [2] De Ceuster, P. (2025). *Interferometric Sector-Gate Nervous Systems: Orchestrating AGI via ParityFuse*.
- [3] De Ceuster, P. (2025). *The Binary-Quantum Bridge for AGI*.