

You don't need attention after all

Decentralized neural consensus, evolved wiring, and brains that learn through experience

Bernhard Mueller

Draft of September 17, 2026

Abstract

We present a novel neural network architecture named Cadence, combining recurrent settling, local equilibrium contrast and sparse record memory for learning while acting. Bounded observer-like patches retain state, communicate through ports and read back consequences, with public evidence of their updates. Fourteen mathematical properties are machine-checked, including conditional convergence, residual error bounds and exact memory identities. A simulated arm reaches held-out targets at 0.994 while learning and at 0.953 with learning frozen; an artist scores 0.708 after a canvas shift with readback and 0.436 without it. Automatic record-bank allocation reaches 0.946 mean accuracy on four synthetic task families after 5,000 decisions against 0.490 for the tested online transformer, with larger memory. A separate solver repairs a localized change on a 138,639-neuron graph 6.30 times faster than full iteration; dense changes favour full iteration. Conventional methods outperform the prototype in several latency, language and music comparisons. We argue for Cadence as an architecture for embodied AGI: robots with bounded local brains that can be taught unfamiliar tasks and retain skills through real-life experience. The proposed efficiency comes from reusing state and updating local computation and memory. The experiments support components of this thesis; physical-robot transfer and lower total resource cost are its decisive tests.

Contents

1	Introduction	2
2	The patch net	4
3	What the animal brain has, and how each part works here	12
4	When retained state saves work	15
5	One stream, one life	16
6	Where the wiring comes from	21
7	A worm that reacts and learns	23
8	Drawbacks	26
9	Open challenges	27
10	Conclusion	27
A	Theorem map and the Lean library	32
B	Evidence map	33
C	Experimental details	34
D	Larger sequence and game controllers	35
E	Notation	36

1 Introduction

1.1 From cognitive competence to embodied learning

Reasoning and conversational competence warrant taking the case for general-purpose cognitive AI seriously. An advanced Gemini Deep Think system achieved gold-medal-level performance on the 2025 International Mathematical Olympiad [46]; a preregistered five-minute Turing-test study found persona-prompted GPT-4.5 selected as human 73% of the time [35]. Whether such capabilities qualify as AGI depends on the required breadth, performance and autonomy [51]. Neither result measures the ability to learn a changing body or retain skills through prolonged physical interaction.

Here *embodied AGI* denotes robots that can be taught a broad range of unfamiliar tasks, learn from their consequences and retain and transfer useful skills under bounded compute, memory and energy budgets. Robot learning through demonstrations, corrections and autonomous experience has concrete precedents in vision-language-action systems [62]. Cadence’s architectural thesis concerns the organization and cost of that learning: persistent local computation, directly writable experience and adaptable connections should make continuing adaptation more economical at matched competence.

A frozen feed-forward model evaluates a learned map at deployment. Recurrent learners, online gradient methods and transformers with key/value caches provide other forms of carried state; backpropagation does not require learning to stop [67, 76]. The question here is which state a learner should retain and update when its task is an ongoing interaction. A robot can accumulate a model of its own consequences, consult that model before acting and revise it after an unexpected outcome. The cost of this process includes its representation, memory, search and every learning update.

Cadence tests a particular organization of that process: recurrent neural state, directly writable records and configurable connections between bounded patches. It separates three questions that an accuracy score alone cannot answer. Does the system learn useful behaviour from one stream? Does it preserve and revise what it learned? Does reusing state save work at the required accuracy and latency? The comparisons retain conventional controls because a dictionary, numerical solver or gradient learner can share those advantages.

1.2 OPH foundations and neural influences

Pragma’s physics and consensus research. Cadence’s organizing framework comes from Pragma Research’s Observer-Patch Holography (OPH): bounded, self-reading patches repair disagreements on declared interfaces. The repair and observation-determined normal-form theorems give conditions for consistent global answers independent of repair order and determined by protected observations [56, 57]. Event algebras define the observable distinctions [53]. Section 2.2 gives the neural specialization; its guarantees are independent of OPH’s physical interpretation.

Eliasmith: integrating perception, cognition and action. Eliasmith’s *How to Build a Brain* [20] was an architectural influence on Cadence. Its integrated treatment of perceptual, cognitive and motor models motivates the ambition of a neural system that behaves as a whole. The arm-and-eye demonstration was inspired by examples on Eliasmith’s website.¹

Hopfield: computation by settling. Hopfield’s graded-response networks compute through collective relaxation and associative recall [32]. Cadence uses recurrent settling with a sufficient contraction condition for a unique fixed-drive equilibrium and an error bound from the residual. This condition supports verification; it does not establish biological fidelity or multiple memory attractors under the same drive.

Equilibrium propagation: local credit. Scellier and Bengio derive objective gradients from free and nudged equilibria under stated symmetry and smoothness assumptions [68]. The centred positive/negative estimator follows Laborieux et al. [42]. Cadence applies these local contrasts and measures their state-storage cost and learning behaviour; the perturbation phases must also be computed.

¹See the group’s Semantic Pointer Architecture and Spaun demonstrations.

Marr and Albus: expansion and memory. Cerebellar theories motivate expanded representations with a plastic readout [2, 47]. Here sparse keys address residual record writes. The normalized write can store one value exactly, while key overlap gives an explicit interference law. Sparse activity alone guarantees neither unlimited capacity nor retention across changing worlds.

Self-modeling robots learn body models through action and sensation [11]; Polyworld evolves circuits with local Hebbian learning and metabolic costs [82]; PathNet evolves pathways trained by backpropagation [21]. Cadence’s organizing constraint is a decentralized model of computation: bounded observer-like patches retain state, exchange activations through declared ports and expose their readback. Useful behaviour must be built from these interfaces, with explicit teaching signals and memory writes.

Here *local repair* has three precise meanings. During inference, neuron i updates its state from its equation residual, $r_i = \sum_j W_{ij} \text{act}(v_j) + c_i - v_i$; agreement means satisfying these coupled equations, not making neighbouring values equal. Synaptic learning changes the equations using endpoint contrasts between perturbed equilibria, with reward acting through local eligibility traces. Record learning corrects a witnessed prediction error at the active sparse addresses. Settling with fixed weights does not itself learn. The contribution is a concrete realization of this discipline with neuron-specific certificates, explicit memory-interference laws and convergence-preserving structural operations, tested through continuous interaction. Local relaxation, equilibrium contrast and delta-rule memory are established ingredients [32, 68, 84]; the technical results below specify what their combination permits and guarantees.

1.3 The hypothesis

The hypothesis is that a reusable developmental procedure can produce local brains that acquire, retain and transfer skills through experience across unfamiliar bodies and tasks. Retained state permits reuse between decisions; sparse records permit direct correction of witnessed outcomes; local synaptic rules and configurable wiring offer a route to distributed adaptation. These are candidate mechanisms for efficient embodied AGI, with separate mathematical conditions and empirical tests. A coordinating global workspace is one proposed organization.

A successful test must hold task-specific wiring out of the design process and count sensing, memory, communication, search and every learning phase, including any pretrained components. Sample efficiency, per-decision latency and total energy are distinct measures. The simulated-body experiments establish parts of the experience protocol; physical transfer, durable retention and lower total cost at matched competence determine whether the architectural thesis succeeds.

1.4 Technical contributions

A patch net retains state across decisions and computes through the interfaces in Figure 1. The paper develops four concrete contributions.

- **A certified recurrent core.** For the implemented neuron, an explicit incoming-weight bound gives a unique equilibrium, a residual error certificate, and bounds on the repair required after changed input or lesions (Theorems 2.6, 2.3, 2.4 and 2.5). These connect local parameters to checkable network behaviour.
- **Writable memory with quantified interference.** The sparse record interface supports exact correction at a normalized key with unit learning rate and an exact law for how that write changes other reads. Separation and capacity bounds state when experience can be stored without overwriting another association (Theorems 2.9–2.12).
- **Structural change with a retained certificate.** Pruning, masking and gain reduction preserve contraction; added projections preserve it within a declared incoming-weight budget (Theorem 2.14). This supplies a mathematically admissible search space for recurrent wiring. Empirical genome searches have their own conditions and do not establish useful behaviour for every admitted graph.
- **An executable experience protocol.** Action, consequence readback, record writes and local learning form one continuing process. Candidate drives query records without writing imagined outcomes as witnessed facts. A neuron-by-neuron reference engine, accelerated implementations and source-bound receipts connect the computational rules to the reported applications (§2.10, §5).

Computational choice	Mechanism evaluated here	Theorem	Measured in
Fixed work per input	work follows the change; an unchanged input needs no repair step	2.4	§5.4
Context and cached activations	potentials, traces and records persist across decisions	2.9	§5.3
Attention to a growing history	a record read and write whose cost is independent of how much has happened	2.13, 2.15	§5.3, §4
Append-only context tokens	one-shot write, exact revision, no interference between separated keys	2.9, 2.12	§5.3
Weights freeze before use	one mode; the same synapses change during use	2.7	§5
Backpropagation through stored iterations	each synapse reads its own two endpoints and one broadcast error	2.7	§5.5
Interference on tested streams	a sparse code confines each update to the records the reading touched	2.12, 2.10	§5
Recomputed hypothetical inputs	alternatives are settled from the current state and scored by records	2.4	§2.6
Damage changes the model	displacement after a lesion is bounded; the same local rule repairs it	2.5	§5.6
The architecture is chosen by hand	the wiring is a genome; budgeted mutations preserve convergence	2.14	§6

Table 1: The argument of the paper in one table. Every theorem is machine-checked in Lean (Appendix A); every measurement is read from a receipt whose hash is recorded in Appendix B.

Decentralization here describes state ownership and update dependencies. The software schedules phases, reduces residuals and coordinates record reads and action selection centrally; sparse-key selection also reads a group of cells. The experiments do not demonstrate an asynchronous deployment across machines. The separately evaluated combined learner includes a conventional gradient component. Distributing the complete controller requires an implementation and its own communication and latency measurements.

1.5 What the evidence establishes

The main empirical case has three parts: closed-loop learning and adaptation in the arm and artist (§5), automatic allocation of useful record banks on controlled task families (§6.2), and conditionally cheaper numerical repair after localized changes (§5.4). The mathematical results in Section 2 state the guarantees and their hypotheses; their Lean declarations and measurement receipts are mapped in the appendices. The connectome, game and sequence experiments test other uses and boundaries.

The larger application brains are wired by hand. Autonomous allocation and evolution search within supplied families of mechanisms. The principal memory experiments supply addresses or clocks; conventional lookup and sparse solvers share several advantages. The arm’s records learner is slower per decision than its small multilayer control, and conventional sequence learners lead on text and music. The embodied-AGI argument therefore concerns an architectural direction, supported by particular learning and state-reuse results, with the full criterion of Section 1.3 requiring a physical-robot evaluation.

2 The patch net

2.1 Neurons, synapses, settling

A net has n neurons. Neuron i holds a potential $v_i \in \mathbb{R}$ and publishes an activation $s_i = \text{act}(v_i)$. Synapses carry effective weights W_{ij} from neuron j to neuron i ; in the library the effective weight is the product of a gain, the number of contacts, a learned efficacy and a presynaptic factor, and the paper writes the product as

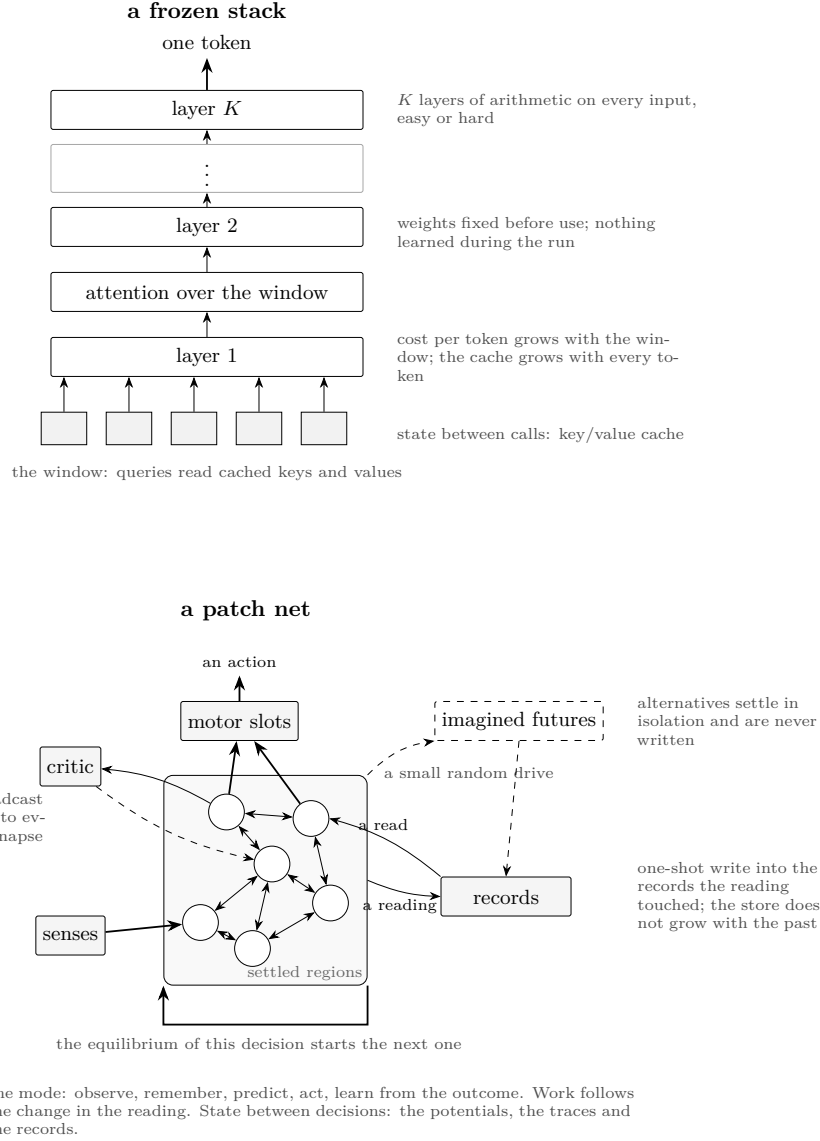


Figure 1: Two organizations of computation. The frozen autoregressive stack shown above applies K layers per new token and retains a growing key/value cache. The patch net below carries its settled state, traces and records into the next decision. Candidate drives are settled and scored using records; witnessed outcomes update those records and local contrast updates the designated synapses. The diagram contrasts these configurations, not every transformer or recurrent learner.

W_{ij} . A stimulus c_i is an additive drive, with the bias absorbed into it. One settling step is

$$v_i \leftarrow (1 - dt) v_i + dt \left(\sum_j W_{ij} \text{act}(v_j) + c_i \right), \quad 0 < dt \leq 1. \quad (1)$$

The activation is a rectified, re-based logistic. With slope λ , threshold θ , leak $\ell \in [0, 1]$ and $\sigma(x) = 1/(1 + e^{-x})$,

$$\rho = \sigma(-\lambda\theta), \quad r(v) = \sigma(\lambda(v - \theta)) - \rho, \quad \text{act}(v) = \frac{\max(r(v), 0)}{1 - \rho} + \ell \frac{\min(r(v), 0)}{\rho}. \quad (2)$$

Rest publishes exactly zero, the activation rises to one, and a neuron below rest publishes a small negative value so that the learning rule can move it. The learning model uses $\lambda = 1$, $\theta = 0$, $\ell = 1/10$. A settling run stops when the largest activation movement in a step falls below a tolerance or when a step cap is reached, and the library reports the residual of equation (1) at the final state. Adaptation, a slow variable subtracted from the drive, produces rhythmic dynamics; the theorems below are stated without it.

The whole brain is one such net. Regions are named ranges of neurons with declared projections between them: senses, association, a workspace that carries context, motor slots, a critic. Settling is joint over all of them, which is what lets a partial reading be completed from the rest of the state.

Intuition: each neuron checks its own account. Imagine an arm whose visual input has moved while its motor and body predictions describe the preceding position. Each neuron compares its held potential with the drive its neighbours and senses supply. Equation (1) moves the potential partway toward that drive; the published activation changes what neighbouring neurons receive, and they make the same comparison. The resulting computation is a feedback process with a measurable discrepancy at every step. In the observer-patch description, each bounded patch owns state, communicates through declared ports, reads back its consequences and retains records. A settling certificate and the public evidence bundle make that self-reading process checkable. The equilibrium reconciles the equations the model implements; whether those equations describe the arm well is a separate empirical question.

2.2 From observable events to a global normal form

OPH distinguishes internal descriptions from shared readings [55–57]. A visual patch may locate an object in camera coordinates and a motor patch in body coordinates. Agreement concerns the position after translation to a shared interface. A change of internal coordinates and corresponding interface map is a gauge change when it preserves every declared public reading. Gauge-invariant repair corrects disagreement in those readings.

For local states x_i and readout maps $\pi_{i,e}$ into a common interface on an overlap $e = \{i, j\}$, agreement means

$$\pi_{i,e}(x_i) = \pi_{j,e}(x_j).$$

A repair preserves the observations held fixed. OPH’s local-repair criterion requires every allowed fixed exterior to admit a consistent interior. A *normal form* has no repair left to make. For a terminating discrete relation on the observable quotient, local confluence gives one normal form from each starting state. Completeness, meaning that normal states coincide with consistent states, makes it a global solution. Local compatibility alone need not imply global consistency. For different interiors with the same protected observations to yield the same answer, the compatible consistent states must also form one equivalence class [56, 57].

What counts as an observable event? For a finite classical record, events are sets of possible readings, closed under union, intersection and complement. OPH’s verified projection-event calculus supplies conditional expectations that discard internal distinctions while preserving the probabilities of the declared partition events [53]. Cadence uses classical activation ports and writable records; the quantum projection and conditioning machinery is mathematical background, without a claimed implementation in this neural model.

The neural specialization. Hold the weights fixed, protect $B(c, v) = c$, and define consistency by $v = W \text{act}(v) + c$. Under Cadence’s contraction condition, every initial state at a fixed input approaches the same equilibrium. The residual and stimulus bounds quantify accuracy and sensitivity. This realizes an observation-determined equilibrium: local equations determine one global answer with a checkable error bound. Numerical

convergence replaces finite exact termination; the synchronous proof does not establish arbitrary asynchronous or transactional repair guarantees. Learning determines whether the answer is useful.

2.3 The certificate: how much work a decision needs

Write $F_c(v)$ for the right-hand side of (1) at stimulus c , $\|\cdot\|_\infty$ for the sup norm on $\mathbb{R}^n$, $\rho_W = \max_i \sum_j |W_{ij}|$ for the largest absolute incoming weight sum (the row mass), and L for a Lipschitz constant of act.

Theorem 2.1 (settling certificate). *If $0 < dt \leq 1$ and $L\rho_W < 1$, then F_c is a contraction of $(\mathbb{R}^n, \|\cdot\|_\infty)$ with constant $q = 1 - dt(1 - L\rho_W) < 1$. [Lean: settle_contracting]*

Corollary 2.2 (unique equilibrium). *Under the hypotheses of Theorem 2.1 the net has exactly one equilibrium v_c^* , and every fixed point of F_c equals it; with zero stimulus and $\text{act}(0) = 0$ that equilibrium is rest. [Lean: equilibrium_unique, silence]*

Theorem 2.3 (a priori and a posteriori certificates). *Under the hypotheses of Theorem 2.1, for every start v and every k ,*

$$\|F_c^k(v) - v_c^*\|_\infty \leq \frac{q^k}{1-q} \|v - F_c(v)\|_\infty, \quad \|F_c^k(v) - v_c^*\|_\infty \leq \frac{\|F_c^k(v) - F_c^{k+1}(v)\|_\infty}{1-q}.$$

[Lean: apriori_bound, aposteriori_bound]

A residual below ϵ at the returned state certifies a potential error of at most $\epsilon/(1-q)$, and an activation error of at most L times that. This is an error bound relative to the model's equilibrium, not a certificate of prediction accuracy. Banach's theorem is the whole proof [6]; what the paper adds is the exact condition on this neuron (Theorem 2.6) and the receipt that the bound never failed in 57,600 checks (§5.4).

Theorem 2.4 (the equilibrium is Lipschitz in the stimulus, and warm starts pay). *Let F_c and $F_{c'}$ satisfy Theorem 2.1 with the same constant q and let $\delta = \|c - c'\|_\infty$. Then*

$$\|v_c^* - v_{c'}^*\|_\infty \leq \frac{dt \delta}{1-q}, \quad \|F_{c'}^k(v_c^*) - v_{c'}^*\|_\infty \leq \frac{dt \delta q^k}{1-q} \quad \text{for every } k.$$

[Lean: equilibrium_lipschitz_stimulus, warm_start_bound]

This is the theorem behind adaptive work. Starting from the old equilibrium, the number of steps that reaches accuracy ϵ after an input change δ is at most the nonnegative ceiling of $\log(dt \delta / ((1-q)\epsilon)) / \log(1/q)$: logarithmic in the size of the change, zero when the change is zero. The bound is sufficient rather than a prediction of semantic difficulty, and checking a returned state has its own cost.

Intuition: reuse the answer, then bound the repair. Contraction says that two possible internal states move closer under the same input: after one step their largest difference is at most q times its previous value. With $q = 1/2$, every step halves that uncertainty. If the residual at a returned state is 10^{-3} , the remaining potential error is at most $2 \cdot 10^{-3}$; the factor $1/(1-q)$ sums all the smaller corrections that can follow. When a sensor changes slightly, the preceding equilibrium is a useful starting point because Theorem 2.4 bounds how far the required answer can move. This can save repeated work on a slowly changing stream. It gives little advantage after a large change, and convergence can be expensive when q is close to one. A unique equilibrium also erases differences in starting state: lasting history must enter through records, traces or a changed drive, rather than through the warm start alone.

Theorem 2.5 (bounded damage). *Let $|\text{act}| \leq A$, let S be a set of neurons, let W^S be W with the columns in S set to zero, and let $\mu_S = \max_i \sum_{j \in S} |W_{ij}|$. If F_c satisfies Theorem 2.1, then the lesioned map satisfies it with the same constant, and $\|v^* - v^{*,S}\|_\infty \leq dt \mu_S A / (1-q)$. [Lean: lesion_bound, lesion_contracting]*

Theorem 2.6 (the library's neuron). *For $\lambda \geq 0$ and $\ell \geq 0$, the activation (2) is Lipschitz with constant $\frac{\lambda}{4} \max(\frac{1}{1-\rho}, \frac{\ell}{\rho})$, satisfies $|\text{act}| \leq \max(1, \ell)$, and publishes exactly zero at rest. Under the learning model the constant is $1/2$ and the bound is 1, so every net whose row mass is below 2 settles by a certified contraction for every $0 < dt \leq 1$. [Lean: activation_lipschitz, activation_abs_le, activation_zero, learningActivation_lipschitz, learning_model_certified]*

For the hypothesis, the certificate makes the numerical state of a changing brain checkable: when its conditions hold, a returned decision carries a bound on its distance from equilibrium. This bound certifies the solve. It does not certify the accuracy of the learned body model or the safety of an action.

2.4 Learning by two equilibria

The free phase settles under the stimulus alone to s^0 . The positive nudged phase settles from s^0 with an extra drive of strength β on the output neurons toward a target: for a quadratic nudge the drive is $\beta(\text{target}_i - s_i)$, for a cross-entropy nudge $\beta(\text{target}_i - p_i)$ with p the softmax of the outputs at temperature T , one softmax per declared slot. The centred estimator settles a negative nudged phase at $-\beta$ from the same free state. With s^+ and s^- the two nudged equilibria, the contrast at synapse (i, j) and at neuron i is

$$\Delta_{ij} = \frac{1}{B} \sum_b \frac{s_{b,i}^+ s_{b,j}^+ - s_{b,i}^- s_{b,j}^-}{2\beta}, \quad \Delta_i = \frac{1}{B} \sum_b \frac{s_{b,i}^+ - s_{b,i}^-}{2\beta}. \quad (3)$$

Reciprocal synapses share one parameter and receive the mean of their two contrasts; a cap bounds every efficacy.

Lemma 2.7 (locality). *There is a fixed function g of four real numbers such that the contrast of synapse (i, j) equals $g(s_i^+, s_j^+, s_i^-, s_j^-)$; the contrast is symmetric in i and j ; and if the two compared equilibria agree, every contrast is zero. [Lean: `contrast_local`, `contrast_symm`, `contrast_eq_zero_of_agree`]*

The lemma describes the information used by the update (Figure 2). Reverse-mode differentiation retains or recomputes intermediate activations; the equilibrium contrast instead reads each synapse’s endpoints in the perturbed states. This can reduce storage when differentiating through many settling iterations. Both perturbed equilibria must be computed, and implicit differentiation offers another way to avoid retaining the iteration history [5].

Fact 2.8 (the contrast is the gradient). *On a smooth stable equilibrium branch with symmetric effective recurrent weights, converged phases and no adaptation, the limit $\beta \rightarrow 0$ of the contrast is minus the gradient of the nudge’s loss with respect to the weight of a reciprocal pair [68]; the centred estimator cancels the first-order bias in β [42]. The library’s tests check the identity numerically: at $\beta = 10^{-4}$ and phase residuals below 10^{-11} the contrast equals the finite-difference gradient to a relative tolerance of 10^{-5} .*

Free and nudged phases are the deterministic form of the two phases of the Boltzmann machine [1, 52], equivalent to backpropagation in a layered net [81], and the local-activation-difference rule of generalized recirculation [59]. Reward uses the three-factor form: an eligibility trace $e \leftarrow \gamma \lambda e + \Delta$ of the same contrast, a linear critic with its own trace, and one broadcast error $\delta = r + \gamma V(s') - V(s)$, with the update $\eta \delta e$ [22, 71, 74].

Intuition: let a small correction travel through the circuit. The target gently pushes the outputs, and recurrence transmits that change to the rest of the net. A synapse compares how strongly its two endpoints coactivate under the positive and negative pushes. Their difference tells it which direction the connection should move under the assumptions of Fact 2.8. Credit reaches the synapse through the circuit’s changed state, so the update itself needs only its endpoints. There is a cost: the learner must solve the perturbed equilibria, and a locally computed update can be slower than a conventional backward pass. An eligibility trace carries that local comparison forward when reward arrives later; it does not guarantee that one delayed scalar identifies the right cause. The advantage is the form and storage of the update, with learning quality established by the controlled comparisons.

2.5 Records: what followed each reading

A record cortex keeps, for each of its cells, what followed the readings that cell took part in. A reading $x \in \mathbb{R}^d$ of declared source neurons is made mean-free by a running mean $\bar{x}$, projected by a fixed random matrix R onto $D \gg d$ cells, rectified, and reduced to the a strongest entries at unit length,

$$k(x) = \text{kwta}_a(\max(R(x - \bar{x}), 0)). \quad (4)$$

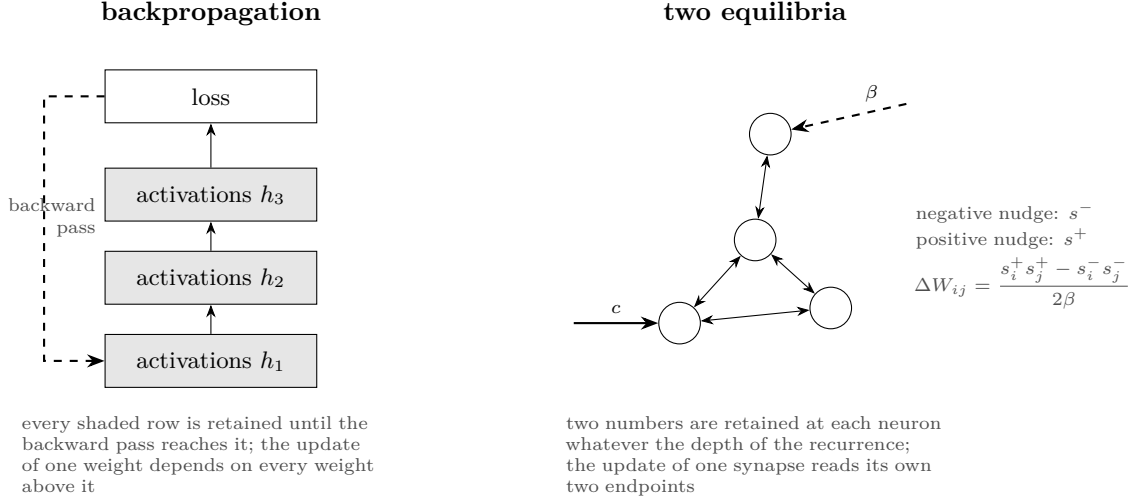


Figure 2: What an update must know. Backpropagation retains the forward tape and propagates a global signal backward through it. The contrast of two equilibria retains two numbers per neuron, and each synapse computes its own change from its two endpoints (Lemma 2.7). The two agree to four digits on the same mechanism (§5.5).

Each predicted field holds one table $M \in \mathbb{R}^{D \times m}$. The prediction is the read $k(x)M$, and the witnessed outcome y is written by the delta rule [38, 79] at rate η ,

$$M \leftarrow M + \eta k(x)^\top (y - k(x)M). \quad (5)$$

Consequence fields use a slow rate so that a regularity is averaged over many outcomes; valued fields such as reward use rate one, so that one exposure writes. Figure 3 shows the parts.

Theorem 2.9 (one-shot memory and its interference law). *For any record M , keys k, q , value y and rate η over a commutative ring: (i) the read at q after the write changes by $\eta(q \cdot k)$ times the correction $y - kM$; (ii) if $k \cdot k = 1$, the read at k after a write at rate one is exactly y ; (iii) if $q \cdot k = 0$, the read at q is unchanged; (iv) if $k \cdot k = 1$, writing the same correct observation twice at rate one changes nothing.* [Lean: read_write, read_write_self, read_write_orth, repeat_write_stable]

Theorem 2.10 (how far one write reaches). *Over the reals, the change a write at k makes to the read at q has sup norm exactly $|\eta| |q \cdot k|$ times the sup norm of the correction, and an overlap of at most ϵ bounds it by $|\eta| \epsilon$ times that norm.* [Lean: read_write_dist, read_write_dist_le]

Theorem 2.11 (exact storage and capacity). *After writing N pairwise orthogonal unit keys in sequence at rate one, every one of them reads its own value, whatever the initial record was; and over a field at most d keys can be pairwise orthogonal unit vectors in a d -dimensional key space.* [Lean: read_writeUpTo, capacity_le]

Theorem 2.12 (pattern separation). *Two keys whose supports are disjoint are orthogonal, so a write at one leaves the read at the other exactly unchanged; a family of unit keys with pairwise disjoint supports, written in sequence at rate one, is stored exactly from any initial record; and a code of width D with a active cells admits at most $\lfloor D/a \rfloor$ such keys.* [Lean: dotProduct_eq_zero_of_disjoint_support, read_write_of_disjoint, read_writeUpTo_of_disjoint, disjoint_family_card_le]

Intuition: a correction follows the overlap of two addresses. Suppose a reading activates two cells equally, giving a unit key $k = (1, 1, 0, 0)/\sqrt{2}$. A write at rate one corrects that reading in a single step. A second reading with key $q = (0, 0, 1, 1)/\sqrt{2}$ shares no cells and keeps its prediction exactly. A third, $q = (0, 1, 1, 0)/\sqrt{2}$, shares one cell and receives half the first reading's correction, since $q \cdot k = 1/2$. This is the interference law in a small example: the same overlap that permits generalisation can overwrite another fact.

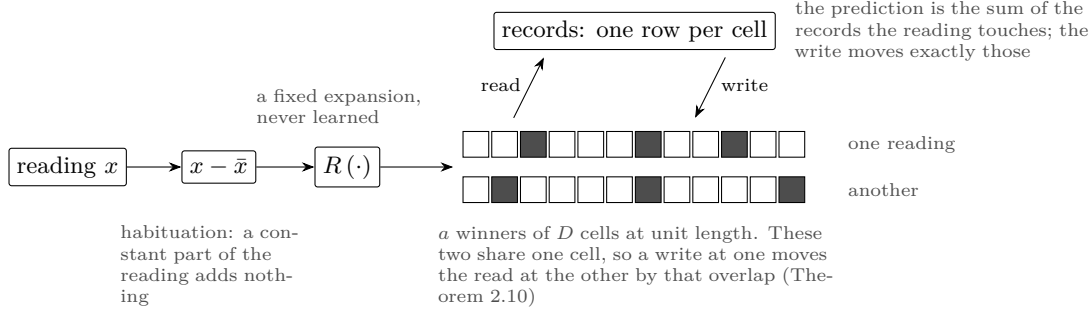


Figure 3: The records cortex. The nonlinearity lives in the code and the learned map is linear and local, the arrangement of the cerebellar granule layer [2, 47] and of the dentate gyrus [48, 60]. Overlap of codes decides interference, and the winner-take-all keeps the overlap small.

The four theorems explain why a useful code can learn quickly from a stream: each observation directly corrects the relevant prediction while disturbing other predictions in proportion to their key overlap. Sparsity alone guarantees neither distinct addresses nor retention; two sparse keys can be identical, and finite capacity limits exact separation. With an appropriate code, direct writes can need fewer observations than fitting the same associations in distributed weights. The record theorem does not by itself prove a replay requirement for another architecture. Section 5 measures that comparison on the declared tasks.

Theorem 2.13 (attention is one read of a record). *Let $M = \sum_t k_t^\top y_t$ be the additive record of N key/value pairs. The read at query q is $qM = \sum_t (q \cdot k_t) y_t$, the stored values weighted by the query/key dot products. If the keys are one-hot at distinct addresses, the read at a one-hot query is exactly the stored value, for every N up to the size of the address space. [Lean: `read_hebb`, `one_hot_recall`]*

The first sentence is unnormalised linear attention [3, 36, 69]; the softmax read of a transformer differs by an exponential kernel and a normaliser, and modern Hopfield networks show that it is one step of an energy network [39, 40, 64]. The residual write of equation (5) is the delta rule over a sequence [84]. The identity covers additive dot-product attention, not the general softmax operation. Unlike an additive write, the residual write corrects its current prediction.

2.6 Symmetry breaking: thinking before acting

A settled state is a fixed point for a particular drive. To consider an alternative, the brain can add a small random drive or encode a candidate action in the input, hold that changed condition during the imagined solve, and settle from the preceding equilibrium. Theorem 2.4 bounds the work after a change in additive drive when the contraction conditions hold. A transient perturbation of state alone returns to the same equilibrium under an unchanged contractive map; it cannot select a second basin. An imagined action produces an imagined reading, whose consequences are read from the records, and a search over such readings chooses what to do. Reading a record changes nothing, so imagination leaves no trace in what the brain has learned; and the imagined state is composed by the same machinery that composes the witnessed one, so a brain can imagine only what its own records support. The arm imagines 364,284 futures per life and the artist about a million, none of which is written [R51,R54]; the Connect Four brain imagines boards through its own learned drop records and searches them with alpha-beta pruning [R53].

2.7 A genome for the wiring

The parts above are fixed within a life. What is open is how they are wired: which fields a record cortex reads, at what gain and offset, how many cells and winners it has, how many pathways each cell samples, which regions project to which, and how many cortices there are. In the library that description is a genome. A genome develops into a brain at a seed, a brain lives one life, and a mutation of the genome changes region sizes, projection densities, signs and gains, adds or removes a cortex, flips a bit of a reading mask, or splits one cortex into two of half the size. Evolution is a loop over genomes with a fitness read from short lives; nothing

acquired in a life passes to the next. The rule that the code is fixed within a life is deliberate: the records are addressed by the code, and a code that changed under a life would misaddress everything written before.

Theorem 2.14 (the certificate under evolution). *Let W satisfy Theorem 2.1 with row mass ρ . (i) Every W' with $|W'_{ij}| \leq |W_{ij}|$ for all i, j satisfies it with the same constant: pruning a synapse, masking a projection, lowering a gain and scaling by any $|g| \leq 1$ are instances. (ii) If V has row mass μ and $L(\rho + \mu) < 1$, then $W + V$ settles by a certified contraction with constant $1 - dt(1 - L(\rho + \mu))$: growth by new projections stays certified inside the row-mass budget. (iii) Every member of a population whose genomes share one row-mass bound settles by one certificate, whatever its individual wiring. [Lean: `rowMassLE_of_abs_le`, `rowMassLE_smul`, `rowMassLE_add`, `mutation_contracting`, `growth_contracting`, `population_contracting`]*

The theorem permits structural search within a checked row-mass budget. It certifies convergence for admitted recurrent weights, not useful behaviour or every empirical genome. Section 6 reports structural searches with their own measurement conditions.

Intuition: search for a useful view of the experience. A record head that reads pitch, gate and unrelated sensory fields together can give similar codes to actions whose consequences differ. Splitting that reading into appropriate heads can make the consequence predictable and protect it from unrelated writes. The genome searches over such choices; experience fills the resulting records. The row-mass budget limits how strongly recurrent connections can amplify one another, so a mutation admitted within that budget preserves convergence. It supplies no test of whether the chosen sensors or divisions are useful. That test comes from held-out predictions or survival under the world’s resource costs. The wiring experiments therefore distinguish finding a useful layout within a supplied family from inventing an unrestricted architecture.

2.8 Cost

Fix the declared cost model. A patch net with E synapses that settles in k steps costs kE multiply-adds per decision, plus Dm for a record read and the same for a write, and holds n potentials and Dm numbers whatever has happened before. A transformer with ℓ layers of width d over a window of T tokens costs about $\ell(12d^2 + 2Td)$ multiply-adds per token and holds a key/value cache of $2\ell Td$ numbers.

Theorem 2.15 (what grows and what does not). *Under that model, the arithmetic a transformer spends on a stream of T tokens is $\ell(12d^2T + dT(T - 1))$, which grows quadratically in T , and its cache grows without bound. The arithmetic a patch net spends is $T(kE + 2Dm)$ and its state is constant in T . For every choice of positive ℓ, d, k, E, D, m there is a horizon beyond which the patch net’s total is the smaller of the two. [Lean: `attention_total`, `attention_total_ge_quadratic`, `record_total`, `record_state_const`, `exists_crossover`]*

The theorem is arithmetic about the declared model, not a claim that the two objects compute the same function. It states where the shapes differ: one of them re-reads its past, the other has written its past into a fixed store. Section 4 puts measured rows beside it.

Intuition: compression trades history for a bounded state. For a repeated body-control task, a learner may need the present position and a learned map from actions to consequences, rather than every observation that established the map. Records and traces carry that summary between decisions. Their fixed size explains the linear stream cost in the theorem, and also imposes a limit: they cannot preserve arbitrarily many independent facts exactly. The quadratic comparison assumes attention to an expanding history; a transformer with a fixed window also has linear total cost in stream length. Whether retaining the summary is better depends on what the task requires, the cost of finding the representation, the achieved error and the measured latency.

2.9 What the theorems tell a designer

Require $L\rho_W < 1$ and $0 < dt \leq 1$ when using the contraction certificate; the learning model has row-mass limit two (Theorem 2.6). Check the potential residual against $\epsilon(1 - q)$ and report both the error bound and the work spent (Theorem 2.3). Warm-start every stream and account for the change in drive (Theorem 2.4). Separate

keys before writing, and measure the overlap after separation: disjoint supports guarantee noninterference, sparse expansion alone reduces it (Theorems 2.10 and 2.12). Give the consequences a slow rate and the reward a fast one. Let a search prune, mask and rescale freely and charge growth against one row-mass budget (Theorem 2.14). Distinguish the lesion bound on potentials from measured recovery under further learning (Theorem 2.5).

2.10 Receipts

Every experiment writes a receipt: canonical JSON whose digest covers the kind, the body and a manifest of source hashes, so that editing any source file invalidates it. A reference engine settles neuron by neuron with one message per declared synapse and a ledger of transmissions, and the accelerated engines are checked against it to $2 \cdot 10^{-16}$ on the CPU. The applications additionally ship a verifier that recomputes every published number from the event logs, and a browser port of the same brain that is checked decision by decision against the Python one. The receipts this paper cites are copied into its evidence package with their hashes (Appendix B).

3 What the animal brain has, and how each part works here

Biological theories motivate several parts of the design, without making it a biophysical model. This section names the relevant mechanisms and the limited correspondences evaluated here. Table 2 gives the map, and Figure 4 shows two application brains while they work.

One stream, and what carries across a moment. The applications learn during an ongoing interaction. Their retained state has three forms: settled potentials, traces and records. Warm starts can save computation when the drive changes little (Theorem 2.4), but convergence to a unique equilibrium erases the initial state at a fixed drive. Unlike persistent-activity accounts of working memory [28, 77], this contraction mechanism therefore needs traces, records or another change to the effective input to retain history.

Short-term memory: the afterglow and the traces. The second kind of carried state is written by no synapse. An afterglow is a trace of a sensory region that decays geometrically after each event, so a moving object leaves a fading streak in the retina’s copy and a motion is readable from one frame; the player of Appendix D has an afterimage of 16,992 neurons beside a retina of the same size. A working-memory trace is the same operation on a cortex rather than a sense: the composer’s prefrontal region is a trace of its phrase cortex with decay 0.85, and its piece region a slower trace with decay 0.98, both supplying retained stimulus to the next equilibrium and read by the phrase cortex through learned synapses [R68]. Fast synapses are the third form, a delta-rule record written every event and decaying, which is what holds a digit span of 12 from position tags [R5]. These explicit stores need no gradient through time. Learning long dependencies through a recurrence can face vanishing or exploding gradients [9, 30]; that is a credit problem, not an absence of recurrent memory. Transformers also retain information in their context and cache. Table 5 tests the particular alternatives on delayed cues.

Long-term memory: plasticity and records. The animal keeps two stores, a fast one that writes an episode in one exposure and a slow one that accumulates structure over many [41, 50]. Here the fast store is the records cortex: a reading lights a sparse code, and what followed is written into exactly those cells by the delta rule of equation (5), in one shot at rate one for a valued field and at a slow rate for a regularity. Sparse conjunctive coding in the dentate gyrus and CA3 keeps the overlap between two episodes small so that one can be stored without overwriting another [44, 60, 75]; Theorem 2.10 is that argument in numbers and Theorem 2.12 its limit. The slow store is every plastic synapse of the settled cortices, changed by the two-equilibria contrast of equation (3) and by the three-factor rule when a reward error arrives. Both stores are written during use and neither is frozen. Section 8 records that the slow store is weaker on the tested task families.

Attention goes to the repair. A local change can leave much of a recurrent equilibrium nearly unchanged. Theorem 2.4 bounds displacement; it does not supply a scheduler or guarantee sparse arithmetic. Figure 4

In the animal	Described by	The part here	Measured in
Granule layer: a few inputs per cell, Golgi inhibition, a sparse expansion read by one Purkinje cell	Marr [47]; Albus [2]; Litwin-Kumar et al. [45]; Cayco-Gajic and Silver [12]	the fixed random projection with winner-take-all of equation (4)	§5.3, §5
Climbing-fibre error at the cell that made the prediction	Ito [33]; Raymond and Medina [66]	the delta write of equation (5), one error per field per cell	§5.3
Dentate gyrus and CA3: conjunctive codes whose overlap bounds interference	Marr [48]; Treves and Rolls [75]; O'Reilly and McClelland [60]; Leutgeb et al. [44]	pattern separation before the write	Thm 2.12, §5.3
Fast hippocampal store beside slow cortical structure	McClelland et al. [50]; Kumaran et al. [41]	two rates: slow for consequences, one exposure for reward	§5
Persistent activity as working memory	Goldman-Rakic [28]; Wang [77]	the carried equilibrium, the afterglow and the traces of §3	§5.3, Appendix D
Adaptation to a constant input	Barak et al. [7] on coding level	the running mean subtracted before the expansion	§5
Recurrent settling that completes a partial input	Hopfield [31, 32]; Rao and Ballard [65]; Friston [24]; Bastos et al. [8]	the joint equilibrium of the settled regions	§5.4
Error-driven learning from local activation differences	O'Reilly [59]; Ackley et al. [1]; Scellier and Bengio [68]; Laborieux et al. [42]	the free and nudged contrast of equation (3)	§5.5
Dopamine reports a reward prediction error on a trace left by recent coactivity	Schultz et al. [71]; Yagishita et al. [83]; Frémaux and Gerstner [22]; Gerstner et al. [26]	the three-factor rule and its eligibility trace	§5
Place cells and a map that survives a change of goal	O'Keefe and Dostrovsky [58]; Hafting et al. [29]; Stachenfeld et al. [73]; Whittington et al. [78]	records keyed by place; a goal changes the reward field alone	§5
Place sequences that depict paths before the animal moves	Pfeiffer and Foster [61]; Mat-tar and Daw [49]	imagined futures from a broken symmetry, searched and discarded	§2.6, §5
A global workspace that broadcasts to every specialist	Baars [4]; Dehaene et al. [17]; Dehaene and Naccache [16]	the jointly settled workspace in the world and composer applications	Appendix D, §9
Synapses that hold state over many timescales	Fusi et al. [25]; Benna and Fusi [10]; Clopath et al. [13]; Frey and Morris [23]	fast and slow strengths with consolidation	§8
Plasticity that does not switch off	Dohare et al. [18]; Zenke and Laborieux [85]	one mode of operation; no frozen checkpoint	§5
Wiring inherited across generations, knowledge acquired within a life	the order of every nervous system	the genome of §2.7; records reset at birth	§6
Measured wiring diagrams	Cook et al. [14]; Dorkenwald et al. [19]; Knox et al. [37]	a connectome loads as the synapse table of the same net	§7

Table 2: Each part of the design, the biological or algorithmic motivation, and the section that measures it.

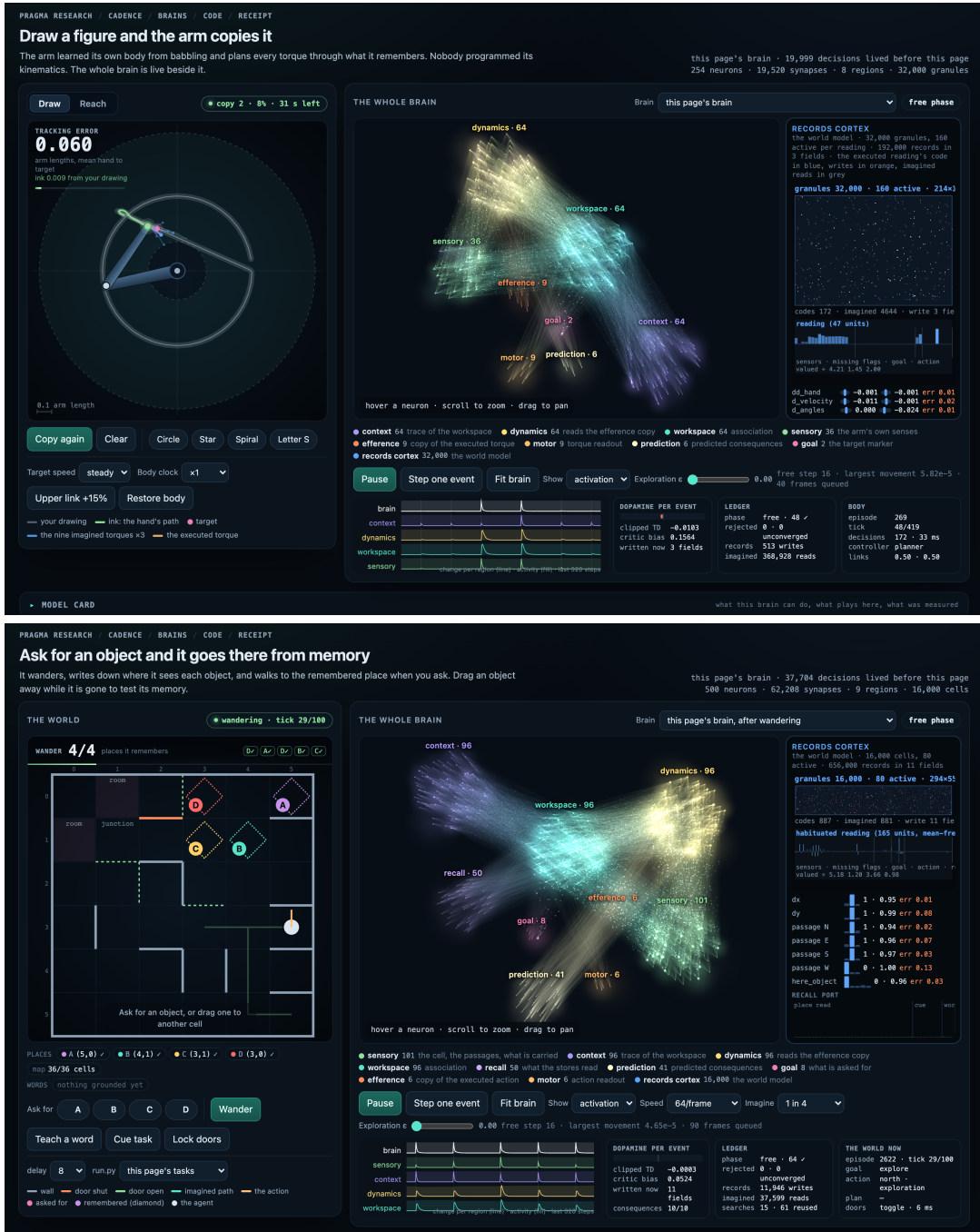


Figure 4: Two of the four lives, captured from the public pages while they run; every panel is the real brain, drawn from its own state. Top: the arm copies a circle a visitor drew, at a tracking error of 0.060 arm lengths, after settling nine imagined torques per decision. The connectome shows eight regions in their colours, brightness for activation, a glow where neurons change, and lit synapses carrying messages; the free phase settled in 16 steps. The right panel is the records cortex, 32,000 granules with 160 active for the executed reading, and the predicted consequences beside their witnessed values, after 368,928 imagined reads and 513 writes. The change per region over the last 320 steps, below the brain, is a display of where activation changed; these traces alone do not show which arithmetic was skipped. Bottom: the world brain wanders a map it has never seen, writes down where it sees each object, and walks there from memory when asked; the ledger shows 11,946 writes and 37,599 imagined reads with no rejected update, and the dopamine panel the clipped reward error of the last event. Both brains keep learning while the visitor plays with them.

displays activation changes. A separate event-driven numerical prototype schedules the affected work and checks the full residual: on the fly graph it answers a local drive change in 5.43 ms against 34.49 ms for full iteration (§5.4). This saving is conditional on locality and scheduling overhead.

Completion, imagination, one scalar for reward, habituation. Recurrent networks that settle into a state which completes a partial input are Hopfield’s construction [31, 32], and predictive coding gives feedback connections the role of carrying a prediction and feedforward connections the role of carrying what it missed [8, 24, 65]; here the joint equilibrium of the settled regions completes a reading with a missing field. Before a rat runs to a remembered goal, place-cell sequences depict the path, including paths never taken [49, 61]; the counterpart is the symmetry-broken settling of Section 2.6. Dopamine neurons report a reward prediction error [71] and convert a recent coactivity flag into a lasting change only within a window of about one second [83], the eligibility trace of three-factor theory [22, 26]; the library broadcasts one such scalar into a local trace at every plastic synapse. Adaptation to a constant input [7] is the running mean subtracted before the expansion, so a constant part of a reading adds nothing to the code.

The expansion and the local readout. A granule cell receives about four mossy-fibre inputs, five orders of magnitude fewer than a Purkinje cell receives parallel fibres, and Golgi feedback keeps the code sparse [45, 47]; Albus [2] cast the circuit as a perceptron whose expansion recoding raises discrimination and learning speed, and the same motif appears in the dentate gyrus [12] and in the fly’s mushroom body [15]. With a fixed expansion and squared-error linear readout the fitting problem is convex, and each update touches only active cells. Stability and retention depend on the step size, key overlap and changing targets. The lives of Section 5 measure whether this memory is useful without replay; they do not prove that another learner must replay.

The global workspace. Global workspace theory holds that a set of specialised processors compete for access to a shared workspace whose contents are broadcast back to all of them, and that this broadcast is what makes a content conscious [4]; Dehaene et al. [17] gave it a neuronal form as a network of long-range neurons that ignites when a content wins the competition, and Dehaene and Naccache [16] set it out as a framework for the neuroscience of conscious access. The hand-wired world and composer brains include a workspace region, a cortex that receives projections from the specialist cortices and projects back, settled jointly with all of them, whose trace is the brain’s context; in the world’s brain of Figure 4 it is the 96-neuron region at the centre. The paper makes no claim about consciousness. It does hold that the workspace is where the hypothesis needs the decision to be made, and Appendix D reports that in the composer the planner reads the records around it.

What the design does not have. Dendritic subunits that compute conjunctions before learning [27, 43, 63] have no counterpart; the neuron is single-compartment. Offline prioritised replay [49, 80] exists in the library as an optional ring and is switched off in every application of Section 5. Metaplastic chains that trade capacity for retention [10, 25] are present only as two timescales.

4 When retained state saves work

A transformer cache avoids recomputing earlier keys and values, although ordinary attention reads a history that grows with context. A fixed-size recurrent state and record bank instead summarize that history at bounded storage cost, trading explicit access to old tokens for finite memory capacity. Theorem 2.15 gives a crossover in its declared arithmetic model, not a universal wall-clock advantage over attention or other bounded-state learners.

Warm starts provide a second possible saving: the required repair shrinks with the input change. A sparse scheduler can exploit localized changes, as the separate solver experiment in Section 5.4 demonstrates. Full residual audits, queue management, initialization and dense changes remain costs. The application brains do not all implement that scheduler.

The learning comparison is likewise conditional. Local contrast can save the tape used to differentiate through settling, but pays for free and perturbed phases. A record can absorb an association in one update,

task	patch net	transformer or best sequence learner
recall over 128 pairs	1.00, 0 trained parameters	0.01, 108,224 parameters, 20,000 steps
revision stream, 128 writes	1.000, 576 bytes per stream	0.145, 17,960 parameters, window of 17,408 bytes
arm, held-out targets, one stream	0.994, no stored transitions	0.402 with 102,790 parameters; 1.000 with a ring of 4,096
world, cue across eight steps	1.000	0.516 with a transformer, 0.588 with a multilayer network
four task families, 5,000 decisions	0.946 with in-life allocation	0.490 online transformer; 0.663 with a ring
four task families, 320,000 decisions	0.981	0.822 online transformer; 0.836 with a ring; 0.866 gradient net with a ring
digit span	12 from 5,696 parameters	12 from 53,770; 5 from 80,138 (LSTM)
three-back, four-back	0.964, 0.954 from 2,744	1.000 from 51,458 with the answer in the window; 0.910, 0.791 from 20,098
cart-pole, steps to threshold	40k from 1,716 parameters, 10.3k multiply-adds per decision	137k from 17,443, 67.8k multiply-adds
decision after an unchanged input	0 repair steps plus a checking step	the full stack
permuted MNIST, mean of thirty	0.913 from 80,339	0.840 from 10,986 (patch transformer)
button calibration	0.455 bits per press	0.806 (GRU), 1.500 (convolutional net)
music, held-out surprise, 65,188 pieces	1.023 from 52.8 million parameters	0.708 from 7.6 million (GRU)
character text, bits per character	3.29 at width 128	3.03 at width 128, 2.67 on validation at epoch 5

Table 3: Every matched row in the paper, wins and losses. Text and music favour the conventional sequence learner. Parameter savings in the memory rows coexist with supplied addresses and with exact dictionary ties (§8).

whereas a gradient learner may need repeated samples. Table 5 measures this on particular streams, counting the controls’ replay; the table also reports the resulting latency disadvantage.

Table 3 collects every row of the paper in which a patch net and a transformer, or the best standard sequence learner, are measured on shared datasets or matched interaction protocols. Arithmetic is not wall-clock, and the same table says where the implementation loses: a patch-net decision on a CPU costs 12.9 ms in the arm’s life against 1.0 ms for a multilayer network in the same life, because settling is sequential and the implementation is a small dense kernel, and on text the conventional sequence learner wins outright.

The receipts report CPU costs of 2,849 seconds for an arm life, 1,933 for the world and about 7,375 for the artist suite, with peak memory of 105 to 115 megabytes; the Connect Four life with its deeper search costs more, at a decision latency of 201 ms at the 95th percentile [R51,R52,R53,R54]. The arm holds 254 neurons with 19,520 directed synapses and 192,000 numbers of records; the world 500 neurons, 62,208 synapses and 656,000 record numbers; Connect Four 187 neurons, 3,960 synapses, 14,000 drop records and 6,000 line records; the artist 511 neurons, 52,160 synapses and 2,240,000 record numbers. No accelerator is used, no corpus is read, and the whole training set is the stream the brain lived through.

5 One stream, one life

5.1 Four lives

Four applications test learning during interaction. Each primary life begins with random synapses and empty records, with no pretraining or replay ring. Frozen evaluations and born-frozen controls separately measure generalization and the contribution of learning. The brain receives moments and returns decisions; the world,

Life	What it learns	What it reaches	Controls that fail
Arm [R51]	the consequences of its own torques, from motor babbling	0.994 of held-out targets reached while learning; 0.953 with learning frozen; copies a visitor’s drawing at tracking error 0.052 against a bar of 0.113; a longer link is absorbed at 0.992 and the original body returns at 0.991	born frozen 0.000, random torques 0.002, shuffled pairing loses 0.224, corrupted dynamics loses 0.404
World [R52]	the consequences of moving, and where it saw what	consequence accuracy 0.979 on held-out episodes; requests from a distant cell fulfilled 1.00; the cue held across eight steps 1.00; a moved object corrected 1.00; doors that stop opening recovered 1.00; words grounded 1.00	memory erased 0.18, words erased 0.10, shuffled pairing gap 0.968, reward-shifted cue 0.514
Connect Four [R53]	what a dropped stone does, which windows are lines, and what its own searches proved, from its own games	exact next boards 1.00 on 10,000 held-out moves, terminal prediction 1.00, tactical suite 1.00, 0.9994 against random and 0.9946 against a one-ply opponent, twelve gates of twelve	corrupted dynamics loses 0.807
Artist [R54]	what its strokes leave on a canvas	held-out foreground 0.932 while learning, 0.885 frozen; a longer link mid-life 0.926 and the original body 0.927	born frozen 0.00, random scribbling 0.08, the supplied path oracle 0.64, an online multilayer network with a replay ring behind the same search 0.80, shifted canvas: 0.708 with readback, 0.436 without

Table 4: The four lives, five acceptance seeds each, every gate passed. None of them writes a replay ring: the receipts record zero replay writes.

	records	MLP	MLP, ring	transformer	transformer, ring
arm, held-out targets reached	0.994	0.018	0.980	0.402	1.000
arm, copying a drawing (error)	0.052	0.702	0.102	0.596	0.058
world, consequence accuracy	0.979	0.949	0.952	0.932	0.948
world, cue across eight steps	1.000	0.588	0.492	0.516	0.484
world, cue across 32 steps	1.000	0.576	0.486	0.494	0.466
parameters or records (arm)	192,000	2,182	2,182	102,790	102,790
decision latency, arm (ms)	12.9	1.0	1.0	5.6	8.5
stored transitions	0	0	4,096	0	4,096

Table 5: Matched life protocols with the world model replaced, five seeds; trajectories depend on each learner’s actions. Held-out evaluations permit continued learning. Without a replay ring the multilayer network reaches 2 percent of the arm’s held-out targets and the transformer 40 percent; with a ring they reach 98 and 100. The records brain reaches 99.4 percent from its own single stream with no ring. On the cue that has to be carried across a delay, every network arm sits near chance whatever its ring. Latency is the price: a decision costs 13 times the multilayer network’s.

its rules and its rewards sit outside it. Receipts bind source hashes, event logs and checkpoints, and verifiers recompute the results. Acceptance uses five fresh seeds at full budget.

5.2 Learning during the life, from one stream

The comparison runners use the same world protocol, sensor encoding, planner, decision budget and five seeds, and replace the consequence learner: the records cortex, an online multilayer network, or an online transformer, each with and without a replay ring of 4,096 stored transitions that supplies one replayed update per real one [R55]. These are separate closed-loop lives: different actions produce different later observations. Table 5 reports adapting held-out evaluations, not frozen generalization on an identical trajectory.

The replay ring substantially improves both conventional models on reaching. The records learner achieves

key cosine	keys	rounds	code cosine	plain record	separated	separated, uncentred	dictionary
0.0	64	4	0.014	0.512	1.000	1.000	1.000
0.5	64	4	0.014	0.325	1.000	0.725	1.000
0.9	16	4	0.014	0.200	1.000	0.500	1.000
0.9	64	4	0.014	0.138	1.000	0.350	1.000

Table 6: Correlated keys under revision, all-key accuracy, five seeds. Sixty-four keys at cosine 0.9 in a 32-dimensional key space, revised four times, are held exactly after separation and at 0.138 without it. The mean absolute cosine between codes is 0.014 whatever the cosine between keys, which is Theorem 2.10 in numbers. Without the running mean the shared component chooses the winners and the record fails. An exact dictionary is perfect here and is the control this row does not beat.

comparable performance without storing transitions, but uses 192,000 record values against the MLP’s 2,182 parameters, and costs 12.9 times its decision latency. The delayed-cue gap also depends on the supplied representation and memory access. The interference identity describes record updates; it is not a theorem about the controls’ replay requirements or their achievable capacity.

The held-out target score permits learning during evaluation and therefore measures adaptation to fresh targets. Freezing the learned arm lowers success from 0.994 to 0.953; the artist’s corresponding foreground score falls from 0.932 to 0.885 [R51,R54]. Both remain competent when frozen, while the adapting scores measure the full experience loop. These comparisons support useful online memory under a shared planner, without a general claim against gradient learning.

5.3 Owned state: one-shot memory and its limits

Memory inside a life. In the world, a request to fetch an object seen earlier from a cell half the map away succeeds 1.00 of the time; with the place records erased right before the request the same histories succeed 0.18 [R52]. When an object is moved in view, the memory is corrected 1.00 of the time. When the doors stop opening, the life recovers 1.00. A reward that follows a cue seen eight steps earlier is chosen 1.00, against 0.514 when the reward is shifted onto the other cue, which is the control that fails when a brain has memorised the room rather than the cue.

Associative recall and revision. A vocabulary of 128 symbols, a context of N key/value pairs, then a query. The net writes each pair as one outer product and answers by settling with the key as stimulus. No parameter is trained. Accuracy is 1.00 at every length from 4 to 128 pairs. A two-layer transformer of 108,224 parameters trained for 20,000 steps on contexts up to 32 pairs reaches 0.26 at four pairs and 0.01 at 128 [R1]; its training loss ended at 3.7 to 4.1 against 4.85 for a uniform prediction, so the row states what that budget buys rather than what attention can do. Under a stream that revises eight keys, the residual record holds 128 writes at 1.000 where a two-layer transformer trained on lengths 8 to 32 falls to 0.145 outside its training lengths and an additive record to 0.260 [R2]. A later token can correct an earlier statement in a context; a residual record instead overwrites its prediction directly (Theorem 2.9(ii)).

Correlated keys. The case where a plain record fails is the one pattern separation is for. Keys of dimension 32 with pairwise cosine up to 0.9, 16 or 64 of them, each revised over four rounds [R3]. Table 6 gives the receipt.

Span and clocked memory. Fast synapses from position tags reach a digit span of 12 with 5,696 parameters where an LSTM trained through time with Adam reaches 5 with 80,138, and a one-layer transformer over the trial reaches 12 with 53,770 [R5]. On n -back, records written from a period- n clock reach 0.943, 0.982, 0.964 and 0.954 at one to four back with 2,744 parameters, against an LSTM’s 1.000, 0.989, 0.910 and 0.791 with 20,098, and a windowed transformer whose window contains the answer at 1.000 with 51,458 [R6]. The position tags and the clock are supplied by the environment, which Section 8 counts as a debt.

change δ	0	10^{-3}	10^{-2}	10^{-1}	$3 \cdot 10^{-1}$	1
warm steps to the new equilibrium	0.0	10.0	14.0	17.3	19.0	20.3
warm steps, the library’s stopping rule	1.0	9.0	13.0	17.0	18.0	20.0
the certificate’s budget	0	27	35	43	47	51
cold steps	21.0	21.0	21.0	21.0	21.0	21.0

Table 7: Compute follows change, row mass 1.0, five seeds. An unchanged input needs zero repair steps and one checking step; a thousandth of a change costs half a cold settling; the count grows with the logarithm of the change. Across the three row masses the a posteriori bound of Theorem 2.3 held in 57,600 of 57,600 checks and the displacement stayed inside the bound of Theorem 2.4 in every cell. A stack of K layers spends K layers in every one of these columns.

rule	new drives	stronger drives	eight lesions	retained state	fit seconds
centred contrast (two equilibria)	4.56865×10^{-8}	1.42443×10^{-7}	1.49987×10^{-7}	110,592 B	0.852
backpropagation through the solver	4.56913×10^{-8}	1.42457×10^{-7}	1.50013×10^{-7}	825,744 B	0.596
implicit gradient	4.56913×10^{-8}	1.42457×10^{-7}	1.50013×10^{-7}	n/a	0.254

Table 8: Three learning rules on one mechanism, three teachers, 256 test cases per condition. The rule that reads two endpoints learns what the tape learns, to four digits, from an eighth of the retained state. The gradient controls fit faster in these implementations.

5.4 Work follows the change

Measured. A layered net under the learning model, efficacies scaled to row masses of 0.5, 1.0 and 1.5, five seeds, 64 stimuli. Each stimulus is settled cold to a tolerance of 10^{-6} ; the input is then changed by a sup-norm amount δ and the settling restarts from the old equilibrium [R8].

Sparse repair on a connectome-sized graph. A separate numerical prototype uses the supplied adult fly graph, 138,639 neurons and 2,700,513 directed synapses, under a known contractive rule [19, 70]. The solver receives the graph and changed source identities; it does not learn them or predict the living animal. Each query starts from a common settled baseline. Across 135 conditions, every timed answer includes a full residual audit [R10]. A single-neuron drive change at tolerance 10^{-7} takes 5.43 ms against 34.49 ms for full iteration, a mean per-seed speedup of 6.30 with 10.4 times fewer synapse visits; 32 changed neurons give 5.84 times and 32 ablations 2.98 times. A change to every neuron makes the queue 2.6 times slower. On digits the settling count is 20 on clean inputs, 20 on noisy ones and 23 on occluded ones [R12]: displacement bounds do not measure semantic difficulty.

This is evidence for conditional numerical reuse on a large sparse graph. Establishing the same benefit in an adaptive agent requires counting initialization, memory, changed-source discovery and scheduling over its actual stream.

5.5 Credit without a tape

A symmetric 48-neuron net with unknown weights, 16 driven inputs, 16 observed outputs, three teachers, output targets only. The centred contrast, backpropagation through 64 tied settling iterations, and an implicit gradient receive identical initial weights, minibatches, Adam arithmetic and stability projection [R13].

That row is Lemma 2.7 and Fact 2.8 on a real net. The consequence is not accuracy. It is that the update can be computed at the synapse while the brain is running, which is what the lives of Table 4 do and what a frozen stack does not.

5.6 Damage, and a body that changes

One pass of 30,000 images; then half the pixels dark for good, or half the hidden neurons removed; then a hundred updates of continued learning under each learner’s own rule [R16]. The patch net goes 0.912 intact, 0.822 blind, 0.899 recovered, and 0.794 lesioned, 0.896 recovered: 0.85 of the loss regained after blindness

	one-ply	minimax	solver 50%	solver 70%	AlphaZero	trap suite
the brain, 1,200 games of one life	0.985	0.805	0.870	0.630		
the brain, deeper search	1.000	0.945	0.935	0.745	0.945 / 0.990	0.895
the brain, threats read off the records	1.000			0.950		0.915
AlphaZero, published recipe	0.738	0.688				
minimax, 1,024 nodes	0.875		0.670	0.395	0.375	0.570
one-ply	0.585		0.585	0.200		0.505

Table 9: Connect Four against graded opponents, paired games with the candidate opening the even games [R56,R60]. The minimax column is the supplied-rules opponent at 1,024 nodes against a matched budget of 1,024 imagined transitions. AlphaZero is the published alpha-zero-general recipe at 24 iterations, met at 25 and at 100 simulations per move over 100 games; that opponent scores 0.738 against one-ply and beats minimax 0.625 to 0.375. The trap suite is 200 positions with a forced line four to ten plies deep, scored by a perfect solver. The threat row is the same frozen records with the threat read switched on, 40 games. Against a perfect opponent moving first every arm scores 0.0: the first player wins Connect Four with perfect play.

where Adam regains 0.50, and 0.87 after the lesion where Adam regains 0.86. The patch net loses more at the moment of damage. Theorem 2.5 bounds the displacement of the equilibrium; recovery is the measured behaviour of the same local rule that learned.

The arm’s lower link is lengthened during the same life: target success is 0.992 with the changed body and 0.991 after restoration [R51]. The artist reaches foreground scores of 0.926 and 0.927 under the corresponding changes [R54]. Connect Four’s score against earlier opponents changes by +0.001 after 200 games against a new mixture [R53]. These are bounded tests of adaptation and retention; no theorem guarantees that all previously correct records survive a new body or task.

5.7 Planning in the head

A patch net can imagine because it has a model it owns: the records answer what would follow a reading, and the settling composes the imagined reading. Connect Four is the clearest case, because the quality of a move can be judged exactly.

The brain sees a board as categorical cells and returns a column. It is given no board object, no rules engine, no terminal oracle and no correct-move labels. It learns from its own games what a dropped stone does, which windows of four cells are completed lines, and, in the current life, which positions its own searches proved and which columns the winners of the games it watched played. It searches over boards it composes itself with alpha-beta pruning inside a budget of imagined transitions, and reads the threats standing on a board off the same line records; the rule that turns threats into a value is supplied, like the search [R53]. After 1,200 games of one life, the exact next board is right on 1.00 of 10,000 held-out moves, terminal prediction is 1.00 and the tactical suite is 1.00 on every seed.

Table 9 evaluates a learned transition model inside supplied alpha-beta search and a supplied threat heuristic. Reading threats improves the score against the 70 percent solver from 0.825 to 0.950 and the trap suite from 0.895 to 0.915 without further learning. These bounded opponents and search budgets do not establish superiority over AlphaZero as a method. The artist offers a more direct feedback test: after a canvas offset its foreground score is 0.708 with readback and 0.436 without it [R54]. Its planner can revise a stroke because the next observation reports what the previous action actually did.

Planning makes a learned consequence model useful for action. The supplied search, candidate budget and heuristic are part of the result and its cost; readback then tests the chosen action against the world. This loop, rather than any single game win, is the capability relevant to the hypothesis.

5.8 A life that does not freeze

Thirty permuted-MNIST tasks, 20,000 images each, seen once, one head, no replay [R17]. The patch net learns the first task at 0.924 against Adam’s 0.920, averages 0.913 against 0.921, and ends the thirtieth at 0.894 against 0.918; a patch transformer averages 0.840 and plain stochastic gradient descent 0.837 [18]. Every learner has forgotten the first task by the end, and the patch net drifts three points where Adam drifts none. That is a loss, recorded in Section 8. Inside a life the same property reads differently: the worm connectome of

family	logged streams		intervention fitness		winner’s wiring	
	evolution	random	evolution	random	accuracy	floor
distractors: one relevant field among eight	+0.540	+0.228	+1.705	+0.386	0.660	0.75
conjunction: two fields jointly	+0.737	+0.380	+1.260	+0.952	0.896	0.80
context: the rule depends on a context unit	+0.291	+0.056	+0.514	+0.164	0.508	0.55
effects: the action decides the outcome	+0.939	+0.856	+1.972	+1.671	1.000	0.95

Table 10: Wiring search on four synthetic task families with a labelled oracle mask, 432 measurements per arm, skill over the joined head on the untouched part of the stream [R61]. Evolution is ahead of the random search on every family under both fitnesses. By accuracy on the logged streams the winner clears the floor on two families; under the intervention fitness the distractor winner predicts the effect of an action at a factual skill of 0.921 and is admitted. On the context family the winner drops the context unit and fails, which is the family where the oracle mask’s own ceiling is 0.560.

Section 7 relearns a reversed odour association to 0.96 or better in every seed [R21], the arm absorbs a changed body and returns [R51], the Connect Four brain keeps its old opponents while learning new ones [R53], and a brain that watched nine recordings of one console game and then learned from its own play improves from 626 to 720 pixels on the trained level and from 372 to 806 on a level it never watched [R26].

6 Where the wiring comes from

The application layouts in Section 5 are hand-wired. In the composer, splitting a joined reading into pitch and gate heads was consequential: neighbouring pitch classes shared 0.567 of their active cells at the original offsets and 0.010 at zero [R65]. Automating such representational choices is a central research question. The following experiments test search over recorded streams, allocation during life and ecological selection, each within a supplied mechanism vocabulary.

6.1 Search over the wiring on a frozen stream

The genome of one record head is a reading mask over the stream’s fields, per-field gains and offsets, the number of cells and winners, a fan-in, and a target field with a delay. The fitness is read from a recorded stream and never from a life: the head’s held-out prediction skill on the required fields, less a charge for cells and synapses, and, under the intervention fitness, the head’s skill at predicting the effect of an action actually executed from a saved state. Two searches run at equal budget, twelve generations of twelve, 432 measurements each: an evolutionary loop with elites and mutation, and a random search over the same prior. The wiring a person found by hand is measured in the same receipt as a disclosed ceiling, and a joined head reading every field is the start every stage writes first.

The composer test carries a substantial prior: intervened target heads must read the swept action field, here pitch. The search can remove other inputs and tune the head; it does not discover the relevance of pitch from an unnamed slot. At a two-block horizon, with 3,510 chip blocks charged for intervention states, the selected evolutionary winner scores +0.368 against random search’s +0.063, the hand-built head’s +0.086 and the joined head’s −0.436 [R61]. Its pitch-only reading uses offsets 0.02, 1,296 cells and 55 winners, found in 251 seconds against 806 for random search. Untouched-stream scores of +0.320 and fresh-stream +0.102 are positive, but both fail the all-field requirement because the sound-presence prediction fails. On a second stream recorded per held note, random search wins, +0.853 against evolution’s +0.438. This is evidence for useful restricted automation, not a general efficiency advantage for evolution.

Transferred into a life, the evolved head improves early pitch accuracy: across four seeds it reaches 0.384 after 6,160 blocks, against 0.103 for the hand-built head. The latter reaches 0.395 at 12,342 blocks and random search’s head reaches 0.397 at 24,926; all end near 0.52 at 50,162 blocks [R65]. The result is faster early acquisition at these checkpoints, not twice the eventual competence. On the 1942 stream, 400 intervention states with eighteen actions each yield no selected wiring better than predicting no effect; the evolutionary winner omits the action and scores 0.000 [R61]. That negative result limits this representation and horizon.

arm	at 5,000	at 320,000	floors cleared at 320,000	lives clearing
in-life allocation, context router	0.946	0.981	4 of 4	8 of 8
in-life allocation, uniform mixture	0.873	0.969	4 of 4	8 of 8
gradient net, encoder pinned at half budget, ring	0.336	0.866	4 of 4	8 of 8
the oracle’s mask (a labelled ceiling)	0.796	0.864	4 of 4	8 of 8
online transformer, ring	0.663	0.836	3 of 4	6 of 8
online gradient net, ring	0.541	0.825	3 of 4	7 of 8
online transformer	0.490	0.822	3 of 4	6 of 8
online gradient net	0.341	0.769	3 of 4	6 of 8
the fixed records brain	0.518	0.551	1 of 4	2 of 8
local contrast, settled regions	0.242	0.202	0 of 4	0 of 8

Table 11: Mean accuracy over the four task families against the budget of decisions in one life [R62]. In-life allocation clears every floor on all forty lives at 5,000 decisions; the online transformer needs sixty-four times the budget to clear three of four and stays at chance on the conjunction family at every budget. The fixed records code is flat after 20,000 decisions and the local-contrast arm falls with budget.

6.2 Allocation during a life, and the budget curves

The allocator starts with a bank reading the full stream. Residual-weighted statistics propose masks from a supplied vocabulary of single coordinates, pairs and the full reading; candidates face probation, with at most six banks and a context router. The policy and mask vocabulary are hand-specified; the selected coordinates are learned. Fourteen arms run on four families at 5,000, 20,000, 80,000 and 320,000 decisions, with ten, ten, three and two lives per family and preregistered floors [R62]. Guards that protect records across world changes are evaluated separately.

The allocator reaches 0.946 mean accuracy after 5,000 decisions and clears every floor on all forty lives, against 0.490 for the online transformer and 0.341 for the gradient net. At 320,000 decisions the transformer reaches 0.822. This is a sample-efficiency comparison, not matched memory: at 5,000 decisions allocator state occupies 1.63–13.33 MB across families, against 0.087–0.112 MB for the online transformer. The fixed records learner plateaus after 20,000 decisions, while local contrast alone falls from 0.242 to 0.202 as the budget grows. More experience therefore does not automatically improve the slow learner.

Acquisition does not imply retention. After learning a second world, the plain allocator recalls the first at 0.599 at the 5,000-decision budget and 0.261 at 320,000. A combined learner adds a conventional gradient patch, a counted replay ring and protected banks. On the 320,000-decision distractor task, the guarded and unguarded versions both acquire at 0.967, while old-world readback is 0.945 with the guard and 0.163 without it, two lives each [R63]. This result belongs to the combination, not local contrast alone. Calibration of a broader developmental policy reaches 0.961, 0.978, 0.945 and 1.000 on four development families but only 0.418 and 0.575 on the delayed-cue and returning-regime transfers [R64]. It does not establish general transfer of a learned developmental rule.

These experiments support automatic selection of useful input views and record allocation within a supplied grammar. They separate acquisition from retention and expose the price of the memory. They provide no example of a mechanism invented outside the genome or allocator vocabulary.

6.3 Ecological selection of modular wiring

Patch World replaces an explicit fitness ranking with reproduction and death in a supplied ecology. A 96×96 torus conserves mass across soil, food and creatures, with drifting light and local interaction rules. A genome specifies up to four record cortices, masks over fourteen sensory groups, sizes, active counts, fan-in, sensing radius, planning horizon and learning rates. Reading, storage and writes cost mass. Energy thresholds govern reproduction; children inherit mutated genomes and empty records. The learner’s internal reward is mass change. Mutation operators explicitly include adding, removing and cost-neutrally splitting a cortex. Thus evolution selects combinations of available parts; neither the parts nor the splitting operation emerge. Runs start with 300 founders and span 6,000–8,000 ticks on two seeds [R67].

The clearest architectural response occurs in rock-and-bite worlds when reading costs a quarter of the base price. The final populations average 3.08 and 2.14 cortices per creature; one contains 167 three-cortex and 126 four-cortex creatures among 361. One-step planning occurs in 201 of 361 and 324 of 513 creatures, with longer

horizons also present. Under the original read price, biting worlds often lose food sensing instead. This contrast supports selection shaped by the cost of information; two seeds and descriptive population comparisons do not isolate each mechanism’s causal contribution.

The receipt’s “halves” statistic means at least two cortices plus a reading mask that includes the prior-code channel. It is 8 and 10 percent in those two cheap-reading worlds and 28 and 15 percent with free reading. The channel makes previous code summaries available for sampling, including a cortex’s own code and unused slots. The statistic checks neither bilateral symmetry nor actual cross-cortex synapses, reciprocity or functional dependence. Median lifetimes of marked versus unmarked creatures are 37 versus 33 ticks in one free-reading seed and 30 versus 31 in the other. There is no consistent longevity advantage. Hearing is associated with reward of 0.50 versus 0.44 per tick, without a causal intervention; symbols carry only 0.001–0.004 bits about the measured food, neighbour and kin variables.

Patch World provides exploratory evidence that ecology and sensing costs can select multi-module record brains from a supplied grammar. It does not demonstrate evolved bilateral organization or a hierarchy of mutually necessary modules. Establishing functional internal readback requires checking the sampled connections and intervening on them; module count alone cannot supply that evidence.

7 A worm that reacts and learns

The *C. elegans* experiment makes the computation visible at the level of named cells. The fixture derived from the hermaphrodite connectome [14] contains 300 connected neurons. It lists 3,638 chemical connection pairs and 1,093 gap-junction pairs. Removing 47 self-pairs, making gap messages bidirectional and merging parallel messages gives 5,044 directed computational links. CANL and CANR have no synaptic edges in this fixture and are omitted. Measured wiring determines the graph; Cadence supplies the neuron equations and fitted parameters. Gap-derived links carry weighted activations, without an electrical conductance model.

7.1 A collision, seen from inside the network

Figure 5 shows the entire graph inside a schematic worm outline and follows one actual boundary encounter from the trained browser model [R70]. The model body is a point controller in a circular dish: it receives odour comparisons and touch, chooses a turn or forward action, and reads back the result. The drawing exposes its internal computation; it does not simulate muscles or the deformation of a living worm.

1. Contact changes the sensory drive. An attempted movement crosses the dish rim. The body code blocks that movement and activates the six supplied FLP/OLQ touch inputs. Their held potentials describe the preceding reading. The changed drive therefore produces a local equation residual, $r_i = \sum_j W_{ij} \text{act}(v_j) + c_i - v_i$.

2. Local updates spread the disturbance. Each neuron moves its potential toward the value its inputs require. Its changed activation alters neighbouring inboxes, so a disturbance initially concentrated near the input can recruit distant cells through recurrent connections. There is no prescribed sequence of layers. In this trace, nine neurons have absolute residual above 0.01 at the changed input, 183 do after four updates, and none do after 35. The solver uses synchronous sweeps; locality describes the dependencies, not a measured saving from skipping neurons.

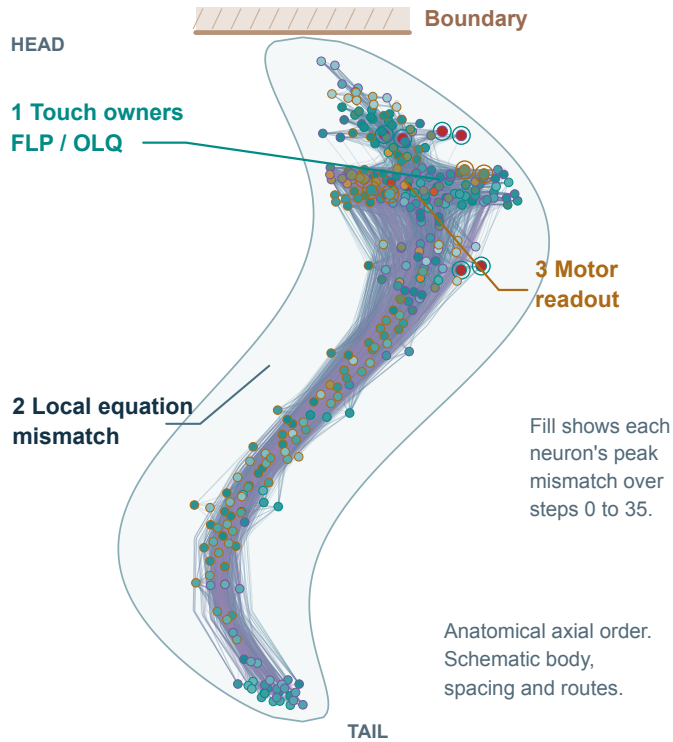
3. The motor readout changes the next action. The supplied output pairs SMDD, SMDV and AVB control left, right and forward choices. With identical post-contact odours, switching touch off gives a right-turn probability of 32.2%; with touch present it is 99.4%. The continued rollout clears the blocked boundary after five actions. This is one controlled illustration of a response, not an obstacle-avoidance benchmark. The trained parameters are frozen throughout the encounter: settling changes neural state, while learning changes the rule governing subsequent encounters.

A boundary encounter through the whole worm graph

Measured *C. elegans* wiring; a recorded response in the Cadence model

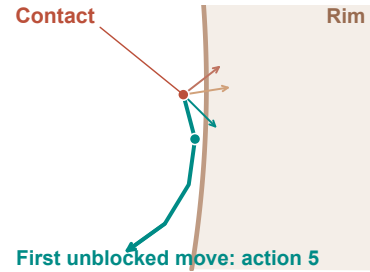
300 modeled neurons / 3,638 chemical pairs / 1,093 gap pairs

A Touch input, local repair, motor readout



B The recorded motion

Point-body trajectory at the rim

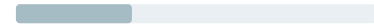


C Touch changes the policy

Right-turn probability

Same odors and starting state

Touch off 32.2%



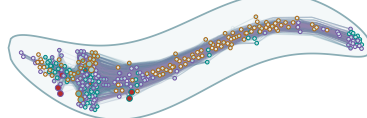
Touch on 99.4%



Frozen weights. One encounter. No new learning in this trace.

D Three actual settling states

Step 0



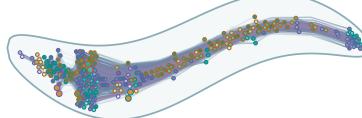
Max mismatch 3.0006

9 neurons above 0.01

Fill: absolute equation residual



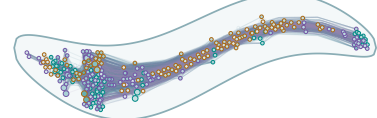
Step 4



Max mismatch 0.3605

183 neurons above 0.01

Step 35



Max mismatch 0.0041

0 neurons above 0.01

Outline: ○ sensory

○ interneuron

○ motor

Figure 5: A disturbance becomes an action on the complete connected worm graph. All 300 modeled neurons and all source connection pairs are drawn; execution uses the 5,044 directed links described in the text. Soma order follows anatomical coordinates, while spacing, body shape and edge routes are schematic. Colours show computed local equation mismatch, with a common scale across the recorded stages. The body trajectory and motor probabilities come from the frozen model, including a control with identical odours and touch removed. Solver steps are not biological time. This figure depicts inference; Table 12 tests learning and reversal.

7.2 What experience changes

In the learning experiment, food is associated with one of two odours, read through AWA and AWC. The agent samples actions, observes outcomes and uses a critic’s reward error to modulate local eligibility traces on 941 designated links; neural biases are frozen. Free and perturbed settling states supply endpoint contrasts. Reward combines food discovery with a supplied approach signal using the simulator’s food location. Thus sensory mismatch drives the state toward a response, while the engineered reward determines which action-related changes to reinforce. The measured graph is directed and asymmetric; applying this learning rule does not establish the symmetric-network gradient identity for the worm.

The three-seed comparison uses sixteen parallel bodies sharing one set of parameters, trained for 8,000 batch updates, or 128,000 body transitions per seed [R21]. With learning disabled at evaluation, the measured wiring finds food in 99.9–100.0% of fresh searches, against 21.3–22.9% before training. Shuffling the source connections yields 52.6–87.7% after the same training budget. The shuffle preserves the raw contact multiset, but merging produces different numbers of computational links. A separate single-body run reaches 100% on 500 evaluation searches after 48,000 decisions; the main three-seed result should therefore not be read as a single 8,000-transition life.

Cells clamped off	Diacetyl rewarded	Butanone rewarded after reversal
None	99.9–100.0%	96.2–99.5%
AWA (diacetyl input)	25.0–35.8%	94.4–99.1%
AWC (butanone input)	99.7–99.8%	24.3–29.4%
RIA (interneurons)	18.9–21.8%	15.9–26.1%

Table 12: **Learning changes which named cells the policy needs.** Food-finding success, range across three seeds, with at least 1,000 fresh searches per seed and intervention, sampled actions and learning disabled. Reversal continues from the acquired parameters for another 8,000 batch updates with the odour–food association exchanged [R21]. The first column comes from acquisition runs; the second from their continued reversal runs. Interventions clamp named cells in the model, including supplied sensory inputs and internal RIA interneurons.

After reversal, AWC becomes necessary for high success and AWA loses most of its influence (Table 12); RIA matters under both associations. This change in intervention dependence connects behavioural learning to named parts of the graph. It provides a more informative mechanism test than activity on the connectome alone. A smaller multilayer actor–critic learns the task faster, and connectome-constrained behavioural models have precedents [34]. The contribution here is the explicit relation between local computation, experience and interventions, without a claim that these equations reproduce the living worm’s physiology. The public learning-worm demonstration exposes the circuit and action loop.

7.3 Other measured connectomes

The adult fly connectome at release 783, 138,639 neurons and 2,700,513 directed synapses [19, 70, 72], settles as one net. Thirty single-channel exposures and 28 held-out mixtures, with localised and distributed lesions of 1,024 and 2,048 neurons, are answered with global residuals below $9.8 \cdot 10^{-9}$ and a uniform potential-error bound of $6.7 \cdot 10^{-8}$, and the library agrees with an independent cold solver to $6.1 \cdot 10^{-16}$ [R22]. Under the published model’s drive convention the untrained wiring passes 14.75 of 29 held-out experimental facts against 7.5 for a shuffled wiring [R23]. The regionalised mesoscale mouse connectome [37], 359 leaf regions per hemisphere and 61,454 projections, is given whiskers, a retina, a nose, a vestibular sense and a place code, with only the three-neuron cord’s 6,462 synapses plastic [R24]: in a labyrinth with food under one of two odours the measured wiring reaches 0.953 held-out success on 64 fresh labyrinths at 1.56 times the shortest path, the shuffled wiring 0.609 at 3.59 times; one seed each, and a torch actor-critic reaches 1.000 in 12 seconds against 46 minutes.

These connectome rows show that the software can instantiate supplied measured connectivity and that this connectivity can matter under the chosen dynamics and sensory interface. They do not establish that those simplified dynamics reproduce a nervous system’s computation. Conventional solvers and other neural implementations can also use the same graph.

8 Drawbacks

1. **The layout is designed by hand.** Every brain of Section 5 and Appendix D was wired by a person, and the wiring was where the time went: the composer’s factorisation cost six rounds and thirty-nine full lives [R65]. Search and evolution have replaced the hand on one record head, on synthetic families, and in a world of small brains; the allocator’s policy parameters are set by hand; the genome configures a mechanism and does not invent one; and the player’s stream defeats the whole search family [R61,R62,R64].
2. **The slow learner.** The settled cortices trained by the local contrast alone fall with budget on the task families, 0.242 to 0.202, and unlearn inside a long life [R62]; the largest brain loses to a recurrent unit on its own corpus, 1.023 against 0.708 [R68]; the composer’s settled regions decide nothing [R65]. The records are a fast memory; the pure local-contrast learner does not turn larger budgets into competence on these families.
3. **Wall-clock.** In the arm’s own life a decision costs 12.9 ms at the 95th percentile against 1.0 ms for the multilayer network [R51,R55]; on supervised rows the factor is 13 to about a hundred [R30,R17,R16,R5]. The player’s conventional arm is 9 to 35 times cheaper per decision [R66], and on the composer the accelerator path is slower than the CPU path [R65].
4. **Language.** On character-level text the patch net is behind at every width tried: 3.29 bits per character at width 128 against 3.22 for a same-window multilayer network and 3.03 for a one-layer transformer whose validation reaches 2.67 [R31,R30]. The joint settling of an utterance is the thesis under test and it is not confirmed [R32,R45].
5. **Hard temporal credit and composed futures.** From reward alone the three-factor rule reaches 174 on Hopper at 300,000 steps against about 500 for proximal policy optimisation, 0.64 to 0.73 of balls on Pong against 0.98, and 488 on a slippery ant at two million steps against 1,373 [R33,R34]. Correcting the policy nudge raises a delayed binary choice from 0.50 to 0.88, where a ten-parameter conventional score rule reaches 0.98 [R44]. Neither reported arcade controller clears a stage [R66,R69].
6. **Retention across worlds.** On split MNIST with one head every learner, this one included, ends between 0.18 and 0.20; a reservoir of 256 stored rows raises it to 0.78 [R36,R42]. On thirty permuted tasks the drift is three points where Adam drifts none [R17,R46]. Read back after a world change, every learning arm of the budget curves loses the old world, and the guard that repairs it for the record banks fails on the effects family above 20,000 decisions [R62,R63].
7. **Supplied addresses, supplied searches, supplied rules.** The position tags of the span task, the clock of n -back, the place keys of the world and the receptive fields of the Connect Four cortices are supplied, and so is every search and the rule that turns threats into a value [R43,R53]. What the brain learns is the model those searches consult.
8. **Conventional controls that win.** A dictionary is exact on every memory row. A windowed transformer is perfect on n -back. A sparse solver settles the fly a tenth faster than the library. A supplied 32-step recurrence answers worm lesions to 0.03 percent [R37]. A multilayer network with a replay ring reaches 0.980 against 0.994 in the arm’s life at a twelfth of the latency [R55], and a gradient net with a pinned encoder and a ring clears every floor of the task families at sixty-four times the budget [R62].
9. **Conditions of the identity and of the certificate.** Asymmetry among free neurons, adapting state and unconverged phases need separate treatment, and finite nudges are biased; a five-seed derivative control measures the residual bias [R41]. Theorem 2.1 is stated without adaptation, and the trained nets of several rows were not checked against its row-mass condition.
10. **Seeds and scale.** The four lives use five acceptance seeds each; the budget curves use two lives at the top budget; the evolution world uses two seeds; historical rows use one to three. The largest brain here holds seventy million synapses, four orders of magnitude below a frontier language model, and nothing in this paper establishes that the design scales to that size.

9 Open challenges

The hypothesis of Section 1.3 depends on several capabilities that the following tests distinguish.

1. **A layout that generalizes beyond its grammar.** Restricted search improves one pitch head, allocation finds predictive coordinates, and ecology selects module counts. Stronger evidence requires learned developmental rules that transfer to new task structures, plus causal tests of the resulting modules. The calibrated policy’s weak delayed-cue and returning-regime transfer [R64] and the player’s zero intervention skill delimit the demonstrated result.
2. **A slow learner beside the fast memory.** The records write a fact in one shot and are flat once their code is used up; the settled cortices are where competence should accumulate, and under the local contrast alone they do not climb. Inside the combined learner a gradient patch does climb, 0.359 to 0.872 [R63]. What is open is a local rule with that property, or a mixed rule whose slow part is local enough to run at the synapse.
3. **A global workspace that decides.** The hand-wired world and composer brains include workspace regions, and in the composer the planner reads the moment and the records while the workspace takes no part [R65]. The hypothesis needs the settled state to be what is acted on, with the records as its memory, which is the order of Section 5’s world and not of the composer.
4. **Retention without a ring.** Sparse writes protect what a reading does not touch; a new world’s readings touch old records anyway, and no policy tested keeps the old world without a guard, and the guard fails where two banks fit two worlds equally [R62,R63]. Metaplastic chains and a reservoir are the known repairs, and both cost memory.
5. **Credit across many decisions and composed futures.** One-step consequences are learned well everywhere; a future composed from them is not trusted by the policy that reads it, and no player clears a stage. The three-factor rule is behind the conventional score rule at long delays [R44,R66].
6. **Bodies and real time.** A robot that learns on the job needs the arm’s property, a body absorbed and returned without a reset, at the latency of its actuators. The arm decides in 12.9 ms and the player in 46 ms at its size and 201 ms at four times its size [R51,R66]; the certificate bounds the error of a decision returned early, and the latency gain from a different substrate must be measured rather than inferred from locality.
7. **Scale and the sequence learner.** On text and on the corpus of the largest brain the conventional sequence learner wins. The hypothesis does not need a patch net to be the better corpus model, but it does need the mechanism behind the composer’s loss understood, and the second item is the current answer.

10 Conclusion

The strongest evidence concerns learning through a feedback loop: records of witnessed consequences support the next action, and a changed body or canvas is read back during the same life. On controlled task families, selecting informative inputs and allocating record banks improves sample efficiency over the tested fixed learners. Under a supplied contractive model, retained numerical state also permits cheaper repairs after localized changes. Each result has a separate evidence contract; none follows from the word “brain.”

The contribution is an explicit local computation and experience protocol, supported by a neuron-specific certificate, quantified memory interference and structural operations that preserve convergence under stated budgets. This makes the consequences of locality and retained state checkable while the learner acts. Wiring searches and Patch World test structural selection inside supplied genome families; they do not establish unrestricted architectural invention or bilateral organization. The theorem that an additive record performs a linear-attention read is an identity about that memory; it is not a demonstration that every use of attention is unnecessary.

We argue for Cadence as an architecture for embodied AGI because its computational organization treats learning from consequences as part of ordinary operation. A robot can draw on pretrained perception or language while learning the particulars of its body and surroundings through continuing interaction. The central hypothesis is that retained local state, writable experience and adaptable wiring can support that learning more economically. Establishing this requires physical robots that acquire unfamiliar tasks, retain earlier skills and transfer across changed conditions at competitive total cost. The prototype supplies mechanisms and bounded evidence for this argument; conventional wins identify where those mechanisms need stronger learning or implementation.

References

- [1] David H. Ackley, Geoffrey E. Hinton, and Terrence J. Sejnowski. A learning algorithm for Boltzmann machines. *Cognitive Science*, 9(1):147–169, 1985. doi: 10.1207/s15516709cog0901_7.
- [2] James S. Albus. A theory of cerebellar function. *Mathematical Biosciences*, 10:25–61, 1971. doi: 10.1016/0025-5564(71)90051-4.
- [3] Jimmy Ba, Geoffrey E. Hinton, Volodymyr Mnih, Joel Z. Leibo, and Catalin Ionescu. Using fast weights to attend to the recent past. In *Advances in Neural Information Processing Systems*, volume 29, 2016. URL <https://arxiv.org/abs/1610.06258>.
- [4] Bernard J. Baars. *A Cognitive Theory of Consciousness*. Cambridge University Press, Cambridge, 1988.
- [5] Shaojie Bai, J. Zico Kolter, and Vladlen Koltun. Deep equilibrium models. In *Advances in Neural Information Processing Systems*, volume 32, 2019. URL <https://arxiv.org/abs/1909.01377>.
- [6] Stefan Banach. Sur les opérations dans les ensembles abstraits et leur application aux équations intégrales. *Fundamenta Mathematicae*, 3:133–181, 1922.
- [7] Omri Barak, Mattia Rigotti, and Stefano Fusi. The sparseness of mixed selectivity neurons controls the generalization-discrimination trade-off. *Journal of Neuroscience*, 33(9):3844–3856, 2013. doi: 10.1523/JNEUROSCI.2753-12.2013.
- [8] Andre M. Bastos, W. Martin Usrey, Rick A. Adams, George R. Mangun, Pascal Fries, and Karl J. Friston. Canonical microcircuits for predictive coding. *Neuron*, 76(4):695–711, 2012. doi: 10.1016/j.neuron.2012.10.038.
- [9] Yoshua Bengio, Patrice Simard, and Paolo Frasconi. Learning long-term dependencies with gradient descent is difficult. *IEEE Transactions on Neural Networks*, 5(2):157–166, 1994.
- [10] Marcus K. Benna and Stefano Fusi. Computational principles of synaptic memory consolidation. *Nature Neuroscience*, 19(12):1697–1706, 2016. doi: 10.1038/nn.4401.
- [11] Josh Bongard, Victor Zykov, and Hod Lipson. Resilient machines through continuous self-modeling. *Science*, 314(5802):1118–1121, 2006. doi: 10.1126/science.1133687.
- [12] N. Alex Cayco-Gajic and R. Angus Silver. Re-evaluating circuit mechanisms underlying pattern separation. *Neuron*, 101(4):584–602, 2019. doi: 10.1016/j.neuron.2019.01.044.
- [13] Claudia Clopath, Lorric Ziegler, Eleni Vasilaki, Lars Büsing, and Wulfram Gerstner. Tag-trigger-consolidation: a model of early and late long-term-potential and depression. *PLoS Computational Biology*, 4(12):e1000248, 2008. doi: 10.1371/journal.pcbi.1000248.
- [14] Steven J. Cook, Travis A. Jarrell, Christopher A. Brittin, Yi Wang, Adam E. Bloniarz, Maksim A. Yakovlev, Ken C. Q. Nguyen, Leo T.-H. Tang, Emily A. Bayer, Janet S. Duerr, Hannes E. Bülow, Oliver Hobert, David H. Hall, and Scott W. Emmons. Whole-animal connectomes of both *caenorhabditis elegans* sexes. *Nature*, 571:63–71, 2019. doi: 10.1038/s41586-019-1352-7.
- [15] Sanjoy Dasgupta, Charles F. Stevens, and Saket Navlakha. A neural algorithm for a fundamental computing problem. *Science*, 358(6364):793–796, 2017. doi: 10.1126/science.aam9868.
- [16] Stanislas Dehaene and Lionel Naccache. Towards a cognitive neuroscience of consciousness: basic evidence and a workspace framework. *Cognition*, 79(1-2):1–37, 2001.
- [17] Stanislas Dehaene, Michel Kerszberg, and Jean-Pierre Changeux. A neuronal model of a global workspace in effortful cognitive tasks. *Proceedings of the National Academy of Sciences*, 95(24):14529–14534, 1998.
- [18] Shibhansh Dohare, J. Fernando Hernandez-Garcia, Qingfeng Lan, Parash Rahman, A. Rupam Mahmood, and Richard S. Sutton. Loss of plasticity in deep continual learning. *Nature*, 632:768–774, 2024. doi: 10.1038/s41586-024-07711-7.

- [19] Sven Dorkenwald, Arie Matsliah, Amy R. Sterling, Philipp Schlegel, Szi-chieh Yu, Claire E. McKellar, Albert Lin, Marta Costa, Katharina Eichler, Yijie Yin, et al. Neuronal wiring diagram of an adult brain. *Nature*, 634:124–138, 2024. doi: 10.1038/s41586-024-07558-y.
- [20] Chris Eliasmith. *How to Build a Brain: A Neural Architecture for Biological Cognition*. Oxford University Press, 2013. ISBN 9780199794546. doi: 10.1093/acprof:oso/9780199794546.001.0001. URL <https://academic.oup.com/book/6263>.
- [21] Chrisantha Fernando, Dylan Banarse, Charles Blundell, Yori Zwols, David Ha, Andrei A. Rusu, Alexander Pritzel, and Daan Wierstra. PathNet: Evolution channels gradient descent in super neural networks, 2017. URL <https://arxiv.org/abs/1701.08734>.
- [22] Nicolas Frémaux and Wulfram Gerstner. Neuromodulated spike-timing-dependent plasticity, and theory of three-factor learning rules. *Frontiers in Neural Circuits*, 9:85, 2016. doi: 10.3389/fncir.2015.00085.
- [23] Uwe Frey and Richard G. M. Morris. Synaptic tagging and long-term potentiation. *Nature*, 385:533–536, 1997. doi: 10.1038/385533a0.
- [24] Karl Friston. A theory of cortical responses. *Philosophical Transactions of the Royal Society B*, 360(1456):815–836, 2005. doi: 10.1098/rstb.2005.1622.
- [25] Stefano Fusi, Patrick J. Drew, and L. F. Abbott. Cascade models of synaptically stored memories. *Neuron*, 45(4):599–611, 2005. doi: 10.1016/j.neuron.2005.02.001.
- [26] Wulfram Gerstner, Marco Lehmann, Vasiliki Liakoni, Dane Corneil, and Johanni Brea. Eligibility traces and plasticity on behavioral time scales: experimental support of neohebbian three-factor learning rules. *Frontiers in Neural Circuits*, 12:53, 2018. doi: 10.3389/fncir.2018.00053.
- [27] Albert Gidon, Timothy Adam Zolnik, Pawel Fidzinski, Felix Bolduan, Athanasia Papoutsis, Panayiota Poirazi, Martin Holtkamp, Imre Vida, and Matthew E. Larkum. Dendritic action potentials and computation in human layer 2/3 cortical neurons. *Science*, 367(6473):83–87, 2020. doi: 10.1126/science.aax6239.
- [28] Patricia S. Goldman-Rakic. Cellular basis of working memory. *Neuron*, 14(3):477–485, 1995.
- [29] Torkel Hafting, Marianne Fyhn, Sturla Molden, May-Britt Moser, and Edvard I. Moser. Microstructure of a spatial map in the entorhinal cortex. *Nature*, 436:801–806, 2005. doi: 10.1038/nature03721.
- [30] Sepp Hochreiter and Jürgen Schmidhuber. Long short-term memory. *Neural Computation*, 9(8):1735–1780, 1997.
- [31] J. J. Hopfield. Neural networks and physical systems with emergent collective computational abilities. *Proceedings of the National Academy of Sciences*, 79(8):2554–2558, 1982. doi: 10.1073/pnas.79.8.2554.
- [32] J. J. Hopfield. Neurons with graded response have collective computational properties like those of two-state neurons. *Proceedings of the National Academy of Sciences*, 81(10):3088–3092, 1984. doi: 10.1073/pnas.81.10.3088.
- [33] Masao Ito. Cerebellar long-term depression: characterization, signal transduction, and functional roles. *Physiological Reviews*, 81(3):1143–1195, 2001. doi: 10.1152/physrev.2001.81.3.1143.
- [34] Eduardo J. Izquierdo and Randall D. Beer. Connecting a connectome to behavior: An ensemble of neuroanatomical models of *C. elegans* klinotaxis. *PLOS Computational Biology*, 9(2):e1002890, 2013. doi: 10.1371/journal.pcbi.1002890.
- [35] Cameron R. Jones and Benjamin K. Bergen. Large language models pass the Turing test. *arXiv preprint arXiv:2503.23674*, 2025. doi: 10.48550/arXiv.2503.23674. URL <https://arxiv.org/abs/2503.23674>.
- [36] Angelos Katharopoulos, Apoorv Vyas, Nikolaos Pappas, and François Fleuret. Transformers are RNNs: fast autoregressive transformers with linear attention. In *International Conference on Machine Learning*, 2020. URL <https://arxiv.org/abs/2006.16236>.
- [37] Joseph E. Knox, Kameron Decker Harris, Nile Graddis, Jennifer D. Whitesell, Hongkui Zeng, Julie A. Harris, Eric Shea-Brown, and Stefan Mihalas. High-resolution data-driven model of the mouse connectome. *Network Neuroscience*, 3(1):217–236, 2019. doi: 10.1162/netn_a_00066.
- [38] Teuvo Kohonen. Correlation matrix memories. *IEEE Transactions on Computers*, C-21(4):353–359, 1972. doi: 10.1109/TC.1972.5008975.
- [39] Dmitry Krotov. A new frontier for Hopfield networks. *Nature Reviews Physics*, 5:366–367, 2023. doi: 10.1038/s42254-023-00595-y.
- [40] Dmitry Krotov and John J. Hopfield. Dense associative memory for pattern recognition. In *Advances in Neural Information Processing Systems*, volume 29, 2016. URL <https://arxiv.org/abs/1606.01164>.
- [41] Dharshan Kumaran, Demis Hassabis, and James L. McClelland. What learning systems do intelligent agents need? complementary learning systems theory updated. *Trends in Cognitive Sciences*, 20(7):512–534, 2016. doi: 10.1016/j.tics.2016.05.004.

- [42] Axel Laborieux, Maxence Ernout, Benjamin Scellier, Yoshua Bengio, Julie Grollier, and Damien Querlioz. Scaling equilibrium propagation to deep ConvNets by drastically reducing its gradient estimator bias. *Frontiers in Neuroscience*, 15: 633674, 2021. doi: 10.3389/fnins.2021.633674.
- [43] Matthew Larkum. A cellular mechanism for cortical associations: an organizing principle for the cerebral cortex. *Trends in Neurosciences*, 36(3):141–151, 2013. doi: 10.1016/j.tins.2012.11.006.
- [44] Jill K. Leutgeb, Stefan Leutgeb, May-Britt Moser, and Edvard I. Moser. Pattern separation in the dentate gyrus and ca3 of the hippocampus. *Science*, 315(5814):961–966, 2007. doi: 10.1126/science.1135801.
- [45] Ashok Litwin-Kumar, Kameron Decker Harris, Richard Axel, Haim Sompolinsky, and L. F. Abbott. Optimal degrees of synaptic connectivity. *Neuron*, 93(5):1153–1164, 2017. doi: 10.1016/j.neuron.2017.01.030.
- [46] Thang Luong and Edward Lockhart. Advanced version of Gemini with Deep Think officially achieves gold-medal standard at the International Mathematical Olympiad. Google DeepMind, 2025. URL <https://deepmind.google/blog/advanced-version-of-gemini-with-deep-think-officially-achieves-gold-medal-standard-at-the-international-mathematical-olympiad/>.
- [47] David Marr. A theory of cerebellar cortex. *The Journal of Physiology*, 202(2):437–470, 1969. doi: 10.1113/jphysiol.1969.sp008820.
- [48] David Marr. Simple memory: a theory for archicortex. *Philosophical Transactions of the Royal Society of London B*, 262(841):23–81, 1971. doi: 10.1098/rstb.1971.0078.
- [49] Marcelo G. Mattar and Nathaniel D. Daw. Prioritized memory access explains planning and hippocampal replay. *Nature Neuroscience*, 21(11):1609–1617, 2018. doi: 10.1038/s41593-018-0232-z.
- [50] James L. McClelland, Bruce L. McNaughton, and Randall C. O’Reilly. Why there are complementary learning systems in the hippocampus and neocortex. *Psychological Review*, 102(3):419–457, 1995. doi: 10.1037/0033-295X.102.3.419.
- [51] Meredith Ringel Morris, Jascha Sohl-Dickstein, Noah Fiedel, Tris Warkentin, Allan Dafoe, Aleksandra Faust, Clement Farabet, and Shane Legg. Position: Levels of AGI for operationalizing progress on the path to AGI. In *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 36308–36321. PMLR, 2024. URL <https://proceedings.mlr.press/v235/morris24b.html>.
- [52] Javier R. Movellan. Contrastive Hebbian learning in the continuous Hopfield model. *Connectionist Models: Proceedings of the 1990 Summer School*, pages 10–17, 1991. doi: 10.1016/B978-1-4832-1448-1.50007-X.
- [53] Bernhard Mueller. Verified projection-event calculus in Lean 4: Bundled arbitrary-partition pinching, Lüders retractions, and CHSH interoperability, 2026. URL <https://floatingpragma.io/oph/papers/machine-checked-finite-event-algebras/>.
- [54] Bernhard Mueller. Observer patch nets: local repair, retained state and the console evidence, 2026. URL <https://github.com/muellerberndt/reverse-engineering-reality>. Observer patch holography corpus.
- [55] Bernhard Mueller. Thinking as patch-net fixed-point search, 2026. URL <https://github.com/muellerberndt/reverse-engineering-reality>. Observer patch holography corpus.
- [56] Bernhard Mueller, Jinwook Kim, David Matscheko, and Jonathan Hill. Observation-determined normal forms: Stability, obstructions, and refinement in constraint and rewrite systems, 2026. URL <https://floatingpragma.io/oph/papers/observation-determined-normal-forms/>. Observer patch holography corpus, machine-checked in Lean 4 with Mathlib.
- [57] Bernhard Mueller, Kai Xue, Jinwook Kim, Arnav Anirudha Kale, David Matscheko, and Jonathan Hill. Reality as a consensus protocol: The fixed-point computation that implements physics, 2026. URL <https://floatingpragma.io/oph/papers/reality-as-consensus-protocol/>.
- [58] John O’Keefe and Jonathan Dostrovsky. The hippocampus as a spatial map. preliminary evidence from unit activity in the freely-moving rat. *Brain Research*, 34(1):171–175, 1971. doi: 10.1016/0006-8993(71)90358-1.
- [59] Randall C. O’Reilly. Biologically plausible error-driven learning using local activation differences: the generalized recirculation algorithm. *Neural Computation*, 8(5):895–938, 1996. doi: 10.1162/neco.1996.8.5.895.
- [60] Randall C. O’Reilly and James L. McClelland. Hippocampal conjunctive encoding, storage, and recall: avoiding a trade-off. *Hippocampus*, 4(6):661–682, 1994. doi: 10.1002/hipo.450040605.
- [61] Brad E. Pfeiffer and David J. Foster. Hippocampal place-cell sequences depict future paths to remembered goals. *Nature*, 497:74–79, 2013. doi: 10.1038/nature12112.
- [62] Physical Intelligence. A VLA that learns from experience. Research report accompanying the RECAP robot-learning study, 2025. URL <https://www.pi.website/blog/pistar06>.
- [63] Panayiota Poirazi, Terrence Brannon, and Bartlett W. Mel. Pyramidal neuron as two-layer neural network. *Neuron*, 37(6): 989–999, 2003. doi: 10.1016/s0896-6273(03)00149-1.

- [64] Hubert Ramsauer, Bernhard Schäfl, Johannes Lehner, Philipp Seidl, Michael Widrich, Thomas Adler, Lukas Gruber, Markus Holzleitner, Milena Pavlović, Geir Kjetil Sandve, Victor Greiff, David Kreil, Michael Kopp, Günter Klambauer, Johannes Brandstetter, and Sepp Hochreiter. Hopfield networks is all you need. In *International Conference on Learning Representations*, 2021. URL <https://arxiv.org/abs/2008.02217>.
- [65] Rajesh P. N. Rao and Dana H. Ballard. Predictive coding in the visual cortex: a functional interpretation of some extra-classical receptive-field effects. *Nature Neuroscience*, 2:79–87, 1999. doi: 10.1038/4580.
- [66] Jennifer L. Raymond and Javier F. Medina. Computational principles of supervised learning in the cerebellum. *Annual Review of Neuroscience*, 41:233–253, 2018. doi: 10.1146/annurev-neuro-080317-061948.
- [67] David E. Rumelhart, Geoffrey E. Hinton, and Ronald J. Williams. Learning representations by back-propagating errors. *Nature*, 323:533–536, 1986. doi: 10.1038/323533a0.
- [68] Benjamin Scellier and Yoshua Bengio. Equilibrium propagation: bridging the gap between energy-based models and back-propagation. *Frontiers in Computational Neuroscience*, 11:24, 2017. doi: 10.3389/fncom.2017.00024.
- [69] Imanol Schlag, Kazuki Irie, and Jürgen Schmidhuber. Linear transformers are secretly fast weight programmers. In *International Conference on Machine Learning*, 2021. URL <https://arxiv.org/abs/2102.11174>.
- [70] Philipp Schlegel, Yijie Yin, Alexander S. Bates, Sven Dorkenwald, Katharina Eichler, Paul Brooks, Daniel S. Han, Marina Yasyukevich, Stefan Berg, et al. Whole-brain annotation and multi-connectome cell typing of *drosophila*. *Nature*, 634: 139–152, 2024. doi: 10.1038/s41586-024-07686-5.
- [71] Wolfram Schultz, Peter Dayan, and P. Read Montague. A neural substrate of prediction and reward. *Science*, 275(5306): 1593–1599, 1997. doi: 10.1126/science.275.5306.1593.
- [72] Philip K. Shiu, Gabriella R. Sterne, Nico Spiller, Anthony Blanć, Zydrune Kliesmete, et al. A *drosophila* computational brain model reveals sensorimotor processing. *Nature*, 634:210–219, 2024. doi: 10.1038/s41586-024-07763-9.
- [73] Kimberly L. Stachenfeld, Matthew M. Botvinick, and Samuel J. Gershman. The hippocampus as a predictive map. *Nature Neuroscience*, 20(11):1643–1653, 2017. doi: 10.1038/nn.4650.
- [74] Richard S. Sutton and Andrew G. Barto. *Reinforcement Learning: An Introduction*. MIT Press, 2 edition, 2018.
- [75] Alessandro Treves and Edmund T. Rolls. Computational analysis of the role of the hippocampus in memory. *Hippocampus*, 4(3):374–391, 1994. doi: 10.1002/hipo.450040319.
- [76] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *Advances in Neural Information Processing Systems*, volume 30, 2017. URL <https://arxiv.org/abs/1706.03762>.
- [77] Xiao-Jing Wang. Synaptic reverberation underlying mnemonic persistent activity. *Trends in Neurosciences*, 24(8):455–463, 2001.
- [78] James C. R. Whittington, Timothy H. Muller, Shirley Mark, Guifen Chen, Caswell Barry, Neil Burgess, and Timothy E. J. Behrens. The tolmán-eichenbaum machine: unifying space and relational memory through generalization in the hippocampal formation. *Cell*, 183(5):1249–1263, 2020. doi: 10.1016/j.cell.2020.10.024.
- [79] Bernard Widrow and Marcian E. Hoff. Adaptive switching circuits. *IRE WESCON Convention Record*, 4:96–104, 1960.
- [80] Matthew A. Wilson and Bruce L. McNaughton. Reactivation of hippocampal ensemble memories during sleep. *Science*, 265(5172):676–679, 1994. doi: 10.1126/science.8036517.
- [81] Xiaohui Xie and H. Sebastian Seung. Equivalence of backpropagation and contrastive Hebbian learning in a layered network. *Neural Computation*, 15(2):441–454, 2003. doi: 10.1162/089976603762552988.
- [82] Larry S. Yaeger. Computational genetics, physiology, metabolism, neural systems, learning, vision, and behavior or PolyWorld: Life in a new context. In Christopher G. Langton, editor, *Artificial Life III*, pages 263–298. Addison-Wesley, 1994. URL <https://shinyverse.org/larry/Yaeger.ALife3.pdf>.
- [83] Sho Yagishita, Akiko Hayashi-Takagi, Graham C. R. Ellis-Davies, Hidetoshi Urakubo, Shin Ishii, and Haruo Kasai. A critical time window for dopamine actions on the structural plasticity of dendritic spines. *Science*, 345(6204):1616–1620, 2014. doi: 10.1126/science.1255514.
- [84] Songlin Yang, Bailin Wang, Yu Zhang, Yikang Shen, and Yoon Kim. Parallelizing linear transformers with the delta rule over sequence length. In *Advances in Neural Information Processing Systems*, volume 37, 2024. URL <https://arxiv.org/abs/2406.06484>.
- [85] Friedemann Zenke and Axel Laborieux. Theories of synaptic memory consolidation and intelligent plasticity for continual learning. *arXiv preprint arXiv:2405.16922*, 2024. doi: 10.48550/arXiv.2405.16922.

A Theorem map and the Lean library

Every Theorem, Lemma and Corollary in the paper names the Lean declaration that proves it. The declarations live in one Lean 4 library (toolchain 4.29.1, Mathlib at the matching release) in seven modules; the axiom audit reports only propositional extensionality, choice and quotient soundness, with no `sorry` and no native evaluation.

Table 13: The theorem map. Settling theorems hold for Lipschitz activations; the Activation module instantiates them for this neuron. Memory identities hold over commutative rings, the capacity bound over fields, the cost statements over the naturals.

paper	module and declaration	content
Thm 2.1	Settling, <code>settle_contracting</code>	contraction with constant $1 - dt(1 - L\rho_W)$
Cor 2.2	Settling, <code>equilibrium_unique</code> , <code>silence</code>	any fixed point is the equilibrium; rest is the equilibrium of no input
Thm 2.3	Settling, <code>apriori_bound</code> , <code>aposteriori_bound</code>	the two error certificates
Thm 2.4	Settling, <code>equilibrium_lipschitz_stimulus</code> , <code>warm_start_bound</code>	the equilibrium moves by at most $dt\delta/(1-q)$; k warm steps leave $dt\delta q^k/(1-q)$
Thm 2.5	Settling, <code>lesion_bound</code> , <code>lesion_contracting</code>	a lesion moves the equilibrium by at most $dt\mu_S A/(1-q)$
Thm 2.6	Activation, <code>activation_lipschitz</code> , <code>activation_abs_le</code> , <code>activation_zero</code> , <code>learningActivation_lipschitz</code> , <code>learning_model_certified</code>	the neuron’s Lipschitz constant, range and rest; row mass below two certifies the learning model
Lem 2.7	Credit, <code>contrast_local</code> , <code>contrast_symm</code> , <code>contrast_eq_zero_of_agree</code>	locality, symmetry, zero contrast for identical states
Thm 2.9	Memory, <code>read_write</code> , <code>read_write_self</code> , <code>read_write_orth</code> , <code>repeat_write_stable</code>	interference law, one-shot write, orthogonal invariance, idempotence
Thm 2.10	Sparse, <code>read_write_dist</code> , <code>read_write_dist_le</code>	the exact size of one write’s effect on another read
Thm 2.11	Memory, <code>read_writeUpTo</code> , <code>capacity_le</code>	exact storage of an orthonormal family; at most d such keys
Thm 2.12	Sparse, <code>dotProduct_eq_zero_of_disjoint_support</code> , <code>read_write_of_disjoint</code> , <code>read_writeUpTo_of_disjoint</code> , <code>disjoint_family_card_le</code>	disjoint supports never interfere; a disjoint family is stored exactly; a code of width D with k winners holds at most $\lfloor D/k \rfloor$ of them
Thm 2.13	Memory, <code>read_hebb</code> , <code>one_hot_recall</code>	the additive read is linear attention; one-hot recall is exact
Thm 2.14	Wiring, <code>rowMassLE_of_abs_le</code> , <code>rowMassLE_smul</code> , <code>rowMassLE_add</code> , <code>mutation_contracting</code> , <code>growth_contracting</code> , <code>population_contracting</code>	pruning, masking and rescaling keep the certificate with the same constant; growth keeps it inside the row-mass budget; one bound certifies a whole population
Thm 2.15	Cost, <code>attention_total</code> , <code>attention_total_ge_quadratic</code> , <code>record_total</code> , <code>record_state_const</code> , <code>exists_crossover</code>	the transformer’s arithmetic over a stream is quadratic in its length and its cache unbounded; the patch net’s is linear with constant state; a crossover exists

Three further results are used as facts and are Lean-checked in the observer-patch corpus rather than in this package [54, 56]: local repair to consensus on a federation of patches computes every finite Boolean function (`echosahedral_universality`); a work-conserving scheduler on a ranked repair system terminates without a fairness premise; and the accepted repair steps of a fixed input-independent federation number at most $n(n+1)/2$, order-sharp. No embedding of those discrete programs into this neuron model is supplied.

B Evidence map

Every number in the paper is read from a receipt copied into the evidence package, with the receipt’s own digest where it has one, the library version it binds, and the sha256 of the copied bytes recorded in the package manifest.

Table 14: Receipt identifiers and manifest labels. The library column names the release whose sources the receipt binds; OPN identifies the earlier observer-patch-nets library. The manifest preserves exact source references and copied-file hashes.

id	manifest label	receipt	library
R1	recall	associative recall over 128 pairs, two transformer arms	0.5.0
R2	revision-memory	memory under a revision stream, five arms, three seeds	0.8.1
R3	pattern-separation	correlated keys under revision streams, six arms, five seeds	0.9.0
R5	digit-span	digit span, six arms, fade sweep	0.7.0
R6	n-back	<i>n</i> -back, four arms, two baselines	0.7.0
R8	adaptive-depth	compute follows change: warm steps, budget, bounds, three row masses, five seeds	0.9.0
R10	event-repair	event-driven repair on the fly graph, 135 conditions	0.8.1
R12	surprise	settling steps against entropy	0.5.0
R13	matched-learning	three learning rules on one 48-neuron mechanism	0.8.1
R16	lesion	damage and recovery, three seeds	0.5.0
R17	permuted-mnist	thirty permuted tasks, six arms	0.5.0
R21	worm-learns	the worm learns an odour, three seeds, lesions, swap	0.8.1
R22	fly-counterfactuals	whole-fly counterfactual screening	0.8.1
R23	opn	fly wiring against shuffled, conformance	OPN
R24	mouse	mouse labyrinth, connectome against shuffled and a network	0.7
R26	player	self-play of the many-games brain	0.7
R27	player	sixteen games at once with a dopamine floor	0.8
R30	ladder-v0.5.0	digits, MNIST, Connect Four, Pong, text, sign, chorales, cart-pole, worm	0.4
R31	sizing	text width sweep, per-run files	0.4 to 0.7
R32	author	language pilots at level one	0.7
R33	hopper, pong	control gates, logs	0.5.0
R34	repeat-previous, slippery-ant	partially observed cards; slippery ant	0.7 to 0.8.1
R36	split-mnist	split MNIST, one head	0.5.0
R37	worm-probe	supplied worm mechanism under lesions	0.8.1
R41	ep-inputs	fixed-input derivatives and free/free asymmetry control	0.9.0
R42	rehearsal	split and permuted MNIST at disclosed memory and work budgets	0.9.0
R43	content-prototypes	noisy cues, fixed prototypes, aliases and capacity	0.9.0
R44	policy-credit	delayed choices, loss-policy correction and score control	0.9.0
R45	sequence-readback	causal held-out text and validation grids	0.9.0
R46	drift-diagnostic	thirty tasks and equal-budget reset forks	0.9.0
R51	experience-arm	the arm: five acceptance seeds, twelve gates, six controls	0.9.0
R52	experience-world	the world: five acceptance seeds, eleven gates, eight controls	0.9.0
R53	experience-connect-four	Connect Four: five acceptance seeds, twelve gates, the threat read and the memory of proved positions	0.9.0
R54	experience-artist	the artist: five acceptance seeds, thirteen gates, eight controls	0.9.0
R55	experience-comparisons	the same lives with the world model replaced by a multilayer network or a transformer, with and without a replay ring	0.9.0
R56	solver-bench	Connect Four against graded solver opponents, every move scored by a perfect solver	0.9.0
R57	instrument	the instrument: five acceptance seeds, ten gates	0.9.0
R58	composer-stage	the chip composer: five acceptance seeds, thirteen gated and seven open measures	0.9.0
R60	connect-four-bench	Connect Four against a published AlphaZero recipe, the threat-read ablation, the trap suite	0.9.0
R61	wiring-search	wiring search on four families and on the composer’s streams: evolution, random search, hand-found wirings, intervention states	0.9.0
R62	scaling-curves	fourteen arms on four task families at four budgets, preregistered floors, retention read-back	0.9.0
R63	combined-learner	banks beside a gradient patch with a guard, ten arms, four budgets	0.9.0

id	manifest label	receipt	library
R64	evolved-development	the evolution loop over wiring and development policy: contract and calibration	0.9.0
R65	composer-curves	the chip composer as curves over a life: three wirings, two sizes, two life lengths, four seeds	0.9.0
R66	player-curves	the 1942 records controller and a PPO convolutional net at matched frames, three seeds	0.9.0
R67	world-evolution	the evolution world: five sweeps over rules and prices, two seeds each, censuses and lifetimes	0.9.0
R68	maestro	the 68-million-synapse musician: receipt, model card and evidence sheet with the recurrent control	0.9.0
R69	player-1943	the 1943 imitation player and its practice on 1942, frozen tests on eight starts	0.9.0
R70	worm-contact	pinned browser sources, complete boundary-contact trace and matched touch-off control	0.8.1

C Experimental details

The four lives. Each life runs one process of small dense arithmetic per seed. The arm decides at 10 Hz with the torque held five body steps, plans by a beam over recorded consequences with a clipped velocity objective, and is evaluated on two isolated copies of each checkpoint, one read-only and one adapting with exploration off. The world is a 6 by 6 map with random connectivity, four walls and three doors, sixteen categorical fields with per-field missing flags, four stores with declared keys, and a best-first search over imagined cells. Connect Four uses two record cortices with receptive fields shared across positions, negamax with alpha-beta and iterative deepening inside a budget of imagined transitions, a validator that rejects an imagined board which is not the old board plus one stone, a memory of positions its searches proved and of the columns winners played, and a life of 1,200 games against random, one-ply, snapshot and solver opponents. The artist uses a 32 by 32 canvas, six pens, an eye of two windows around the gaze and the hand, and a two-level search over stroke intentions and the torques that realise them. Acceptance is seeds 10 to 14 at full budget in every case.

Learning configuration. Unless a receipt states otherwise, supervised rows use the learning model of Theorem 2.6 with $dt = 0.5$, a cross-entropy nudge at $\beta = 0.1$, a free phase of up to 100 steps and nudged phases of 12 to 50 steps at tolerance 10^{-4} to $3 \cdot 10^{-3}$. Records use a rate of 0.2 for consequence fields and 1.0 for valued fields and 0.5 to 1 percent of cells active. Baselines are trained in the same script on the same split with Adam, selected on the same validation set, and read the test set once.

Wiring search and budget curves. The four families are streams of 8,000 decisions over 24 opaque coordinates, permuted and rescaled per world seed, with a required field whose rule names one, two, or two plus a context coordinate, or the action; the oracle mask is the labelled ceiling and the floors are 0.75, 0.80, 0.55 and 0.95. The search splits a stream by whole episodes into a fitting half, a selection quarter and an untouched quarter, scores on development seeds 11 to 13 and reads acceptance seeds 21 to 23 once. The budget curves run each arm at 5,000, 20,000, 80,000 and 320,000 decisions on ten, ten, three and two lives, and read the first world back with learning off after 5,000 decisions in a second world. The evolution world runs 300 founders to a cap of 1,500 at 6,000 or 8,000 ticks, seeds 1 and 2, with the read price at 0.0004, 0.0001 or zero per unit.

Worm contact trace. The recorded encounter uses the exported trained browser model, dish seed 7 and action-sampling seed 17, selecting the first boundary contact without conditioning on the response direction. It occurs at decision 65 of the rollout, during the third search. A touch-off control starts from the identical neural state and preserves post-contact odours. All weights and biases are frozen. The 35-step run stops when the maximum activation movement falls below 0.001; the final equation residual is 0.00406. The incoming-weight contraction condition is not established for this fitted graph, so the displayed settling is an observed response rather than a certified unique equilibrium. The source snapshot, per-neuron trace and deterministic replay checks accompany the evidence [R70].

Connectome rows. The worm, fly and mouse wirings load from their public sources with pinned hashes, signs from transmitter annotations where available. The fly row imposes a row mass of 0.85 so that Theorem 2.1 applies and the residual certifies the state. The worm and mouse rows use the library’s own neuron with a gain chosen so that the net is neither silent nor saturated.

D Larger sequence and game controllers

The hypothesis is about complex multi-cortex brains, so the paper reports the largest ones trained, with the results as they stand.

The composer. The largest is a musician of 17,855 neurons in fourteen regions with 68,570,458 directed synapses and 52,849,100 trainable parameters: an ear, a sense, a mood and an interval region projecting into melody, harmony, rhythm and timbre cortices, a form cortex, a phrase cortex, a prefrontal working memory written as a trace, a recall record keyed by the bar of a sixteen-bar cycle, a slow trace of the piece so far, and an intention region of five softmax slots, all settled in one joint equilibrium and learned by the centred contrast with no backpropagation graph [R68]. It imitates a corpus of 65,188 public-domain pieces, 46.8 million events, one event at a time. Its held-out surprise is 1.023 nats per event after 38,000 updates and 5.2 hours on one accelerator; a gated recurrent unit of 7.6 million parameters trained through time on the same streams reaches 0.708 in 269 seconds. Width did not close that gap and neither did settling depth; the settling probe puts the brain’s surprise at equilibrium 0.05 below the 32-step figure and no closer. This is the loss the hypothesis has to explain, and Section 9 names the mechanism.

A second composer is built the other way round, as one life on an emulated sound chip: 2,355 settled neurons, a world head of 4,000 cells, a phrase store, a preference store, a theme store of 16,000 cells, and voice heads reading the pitch slot, with 541 motor neurons over 28 fields and an ear of 68 features per channel [R58]. On five acceptance seeds it passes ten of thirteen gates: its continuation surprise is 1.270 nats against a gate of 1.608, it beats a procedural control by 0.371, a shuffled pairing by 0.389 and its born-frozen twin by 0.407, and it decides within 18.5 ms. It fails the two musical gates, pitch within a semitone at 0.407 and onset agreement at 0.498 against 0.85, where an echo of the demonstration’s own registers reaches 0.916 and 0.924. A theme it is asked to bring back returns at 0.432 and at 0.011 with the theme store erased. Measured as curves over a life, imitation climbs from 0.10 after 6,160 blocks to 0.53 after 50,162 and stops; four times the life changes pitch by -0.037 ± 0.301 and four times the brain by -0.029 ± 0.163 , every interval covering zero; an instruction objective takes playing under a request from 0.260 to 0.620 and is flat from the first point of every life [R65]. The settled regions of that brain, workspace, dynamics, context, recall, goal and motor, take no part in a decision: the planner reads the moment and the records.

The instrument. The chip alone, learned as one life, passes all ten of its gates on five seeds: held-out prediction gain 0.591 over persistence and 0.741 on transitions, prediction retained after the matching phase at +0.021, a median pitch error of 0.0027 cents measured on the audio, silence when nothing is asked 0.991, a born-frozen tone loss of 0.798 and a shuffled-pairing loss of 0.442, at a decision latency of 8.0 ms [R57].

The player. A brain of 38,577 neurons and 1,056,874 synapses, a retina of 16,992 neurons with an afterimage of the same size, a body, an efference copy, a route memory, a context echo, foveal and peripheral feature maps of 1,800 and 1,680 neurons, an association cortex of 512 and a motor region of nine, watches 275,960 frames of a recorded expert run of the console game 1943 and learns to press the pad: 1,500 updates on one accelerator in 332 seconds take the action loss from 1.101 to 0.867 bits per slot [R69]. Frozen and set down on eight untouched starts it scores 6,062 points and survives 601 decisions on the game it watched; on 1942, a game it never watched, practice from its own play takes it from 369 to 906 points. A records controller built for the same game, 1,413 neurons with a record cortex of 16,000 cells, is measured against a convolutional actor-critic trained by proximal policy optimisation on the same view at matched frames: it is ahead by 1,359, 2,925, 475, 2,009 and 183 points at 12,500 to 200,000 decisions, its best checkpoint scores 6,175, and both sit below a scripted sweep at 7,975 to 8,850; neither clears the first stage in 1,800 held-out episodes, and the conventional arm is 9 to 35 times cheaper per decision [R66]. A brain playing sixteen console games at once with a per-stream dopamine takes 1942, which it never watched, from 350 to 2,067 points in 2,048 episodes and gives the platform games back when continued [R27].

For the hypothesis these are the scale results as they stand: a brain of seventy million synapses settles, learns and composes, and a brain of a million synapses watches a person play and then plays; on the measure that a sequence learner is built for, the sequence learner wins, and on behaviour that needs credit across many decisions, nothing here clears a stage.

E Notation

n neurons; v_i potential; $s_i = \text{act}(v_i)$ activation; W_{ij} effective weight from j to i ; c_i stimulus with bias absorbed; dt step; ρ_W largest absolute incoming weight sum; L Lipschitz constant of act; q contraction constant; v^* equilibrium; β nudge strength; $s^\pm$ nudged equilibria; Δ contrast; η learning rate; M record table; k, q keys; y value; D code width; a active cells of a code; $k(x)$ the code of a reading, which is its key; R the expansion; e eligibility trace; δ the broadcast reward error, or the size of a stimulus change where stated; γ, λ discount and trace decay; μ the row mass of added projections; ℓ, d, T layers, width and window of the transformer in the cost model.