

TECHNICAL THESIS AND DEVELOPMENT PROGRAM

Integrated Whole-Brain Modeling Across Modalities, Scales, and Dynamics for Individualized Neurotechnology

A falsifiable scientific position, executable data contract, and implementation roadmap

Jacob Valdez
Founding Architect, SuperCognition Labs · jacob@supercognitionlabs.com

STATUS	DRAFT. PENDING DEVELOPMENT	DOCUMENT	Thesis / position / implementation contract
	No empirical result is asserted	RELEASE	V6
CLAIM STATE			

ABSTRACT

Magnetic resonance imaging, electroencephalography, transcranial magnetic stimulation, and transcranial focused ultrasound neither observe nor perturb the human brain at commensurate spatial or temporal resolutions. Structural MRI constrains anatomy; fMRI samples blurred vascular consequences over seconds; EEG resolves electrical dynamics over milliseconds through an ill-posed forward model; TMS delivers a state-dependent electric field; and tFUS delivers acoustically distorted exposure with device- and skull-dependent uncertainty. This thesis proposes **SuperCognition Whole-Brain Dynamics (SC-WBD)** as a human-adult model family for relating these methods without pretending that they occupy one raster, clock, coordinate frame, or biological level.

SC-WBD treats a region as a structured, scale-indexed state space rather than a scalar node. Cortical fields, laminar tensors, event processes, associative codes, vascular compartments, and reduced control states communicate through typed operators compiled from a probabilistic anatomical schema. Evidence remains authoritative at its native scale. Fine-authoritative, multilevel-consensus, and coarse-with-sparse-refinement policies allow several resolutions to coexist; calibrated transforms move predictions, gradients, and uncertainty across supports without manufacturing fine truth from coarse measurements. Each source, transform, observation, intervention, and learned residual carries provenance, variance, bias bounds, and model-discrepancy assumptions.

The central claim is deliberately falsifiable: joint, multirate inference should identify selected coupling and delay parameters more reliably than isolated modalities or naive resampling, and a calibrated perturbation should recover distinctions that passive measurements leave unresolved. The first required evidence is therefore a three-region linear–Gaussian recovery benchmark with EEG-like fast mixing, fMRI-like delayed hemodynamics, and an impulse intervention, followed by an end-to-end synthetic compiler demonstration. Both have now been run, and they narrow the central claim: what they support is source-native *clock* handling over resampling, whereas in that deliberately small system the second modality supplies roughly one percent of the electrical channel’s information about coupling—consumed by its own observation nuisances—and effectively none about the conduction delay, and the impulse’s apparent advantage is accounted for by its input energy. These are results about a synthetic design, and about a design shaped to one of its two modalities; no empirical result about brains is asserted. The architecture rejects undeclared units or frames, unsupported prolongation, intervention-free causal labels, unconstrained mechanistic residuals, invalid leakage splits, and cross-scale gluing that fails declared commutation tolerances.

This document states the implementation contracts, rejection conditions, experimental gates, and staged training program. Heterogeneous public, partnered, prospective, and simulated trajectories may supervise only the modules and ports they observe; missing states are marginalized rather than labeled. The near-term objective is calibrated cross-method prediction and bounded TMS/tFUS target comparison, not an unrestricted digital brain. Longer-term language, affective, social, and person-specific models remain downstream hypotheses whose credibility depends on prospective neural, physiological, behavioral, and intervention forecasts.

KEYWORDS

whole-brain modeling; heterogeneous data fusion; multiscale neural dynamics; identifiability; uncertainty quantification; individualized neurotechnology; neuromodulation; computational neuroscience

0. Program Thesis, Claim Boundary, and Implementation Contract

This document is not a report that SC-WBD has already achieved whole-brain prediction. It is the technical thesis that governs what SuperCognition Labs will build, what evidence may update it, which configurations the compiler must reject, and which experiments are capable of changing the program’s scientific commitments. Its first audience is therefore dual: researchers should be able to locate the assumptions that require evidence, while implementation agents should be able to translate each contract into schemas, validators, tests, and acceptance gates.

The thesis has four load-bearing commitments. First, heterogeneous methods are related through a latent generative process and method-specific read/write operators, not through naive resampling into a common tensor. Second, adult human anatomy supplies strong probabilistic topology and phenotype priors, while effective coupling, local computation, state, and slow plasticity remain learnable. Third, resolution is an indexed property of state and evidence: several non-equivalent spatial and temporal supports may coexist without a single materialized truth. Fourth, mechanistic language is earned only by predictions that an equal-capacity surrogate misses, especially under held-out perturbation. These commitments are compatible with useful predictive models that remain agnostic about many biological details; they do not license a claim that every admitted operator is neurally realized.

0.1 What SC-WBD Forbids

A type system earns credibility by rejecting programs. SC-WBD therefore distinguishes an operator family that can be represented from a model instance that is admissible for a particular scientific claim. The compiler must fail closed on the following conditions unless the schema explicitly changes the claim class and records the override.

0.2 Position Relative to Existing Model Families

The latent-state/observation-operator decomposition is not new. Dynamic causal modeling (DCM) already formalizes Bayesian inference over hidden neuronal dynamics, effective connectivity, experimental inputs, and modality-specific forward models (Friston et al., 2003). The Virtual Brain provides an individualizable connectome-coupled simulation environment (Sanz Leon et al., 2013; Sanz-Leon et al., 2015); neurolib and BrainPy provide extensible machinery for whole-brain neural-mass and general brain-dynamics programming (Cakan et al., 2023; Wang et al., 2023). NeuroML/LEMS and PyNN establish substantial prior art for declarative, unit-aware, simulator-independent model specification (Cannon et al., 2014; Davison et al., 2009). SC-WBD should interoperate with these ecosystems rather than redescribe them as absent.

The proposed differentiator is the joint compilation of (i) heterogeneous operator-valued regional state; (ii) non-nested, source-native spatial and temporal resolutions; (iii) uncertain coordinate/clock/field transforms; (iv) data-use provenance into gradient permissions and leakage barriers; and (v) explicit claim-dependent refusal rules. DCM remains the clearest statistical ancestor for generative inversion; declarative neuroscience languages are the clearest software ancestor for typed execution. SC-WBD’s burden is to demonstrate that extending those ideas to multirate, multiresolution, intervention-aware fusion yields identifiable and calibrated predictions unavailable to simpler constructions.

Patient-specific cardiac modeling is a useful neighboring

discipline because it has had to separate code verification, physical validation, parameter calibration, model-form uncertainty, and context-of-use credibility before a digital-twin claim can influence a decision (Galappaththige et al., 2022). SC-WBD adopts the same ladder: numerical correctness is necessary; agreement with recorded signals is stronger; held-out perturbation is stronger still; and prospective decision improvement is a distinct claim.

0.3 Identifiability Before Scale

The first scientific artifact is intentionally small. Consider three latent regional states with directed coupling, one unknown conduction delay, and a controlled input:

$$dx(t) = A(\theta)x(t) dt + B u(t - \tau_u) dt + Q^{1/2} dW_t, \quad x(t) \in \mathbb{R}^3. \quad (T1)$$

An EEG-like instrument observes rapid instantaneous mixing,

$$y_E[k] = L(\ell)x(k\Delta_E) + \epsilon_E[k], \quad \Delta_E = 1 \text{ ms}, \quad (T2)$$

whereas an fMRI-like instrument observes a slow hemodynamic convolution,

$$y_B[n] = M \int_0^{T_h} h(s; \rho) x(n\Delta_B - s) ds + \epsilon_B[n], \quad \Delta_B = 1 \text{ s}. \quad (T3)$$

The unknown vector θ contains selected coupling gains and the network delay, while ℓ and ρ contain observation nuisance parameters. For a linear-Gaussian discretization the expected Fisher information for design d is

$$\mathcal{I}_d(\eta) = \sum_{m \in d} J_m(\eta)^\top R_m^{-1} J_m(\eta) + \mathcal{I}_{\text{prior}}, \quad J_m = \frac{\partial \mu_m}{\partial \eta}, \quad (T4)$$

with $\eta = (\theta, \ell, \rho)$.

Two properties of (T4) must be stated where it is introduced, because a benchmark that ignores either of them will report an artifact of its own construction as a finding.

Modality additivity is an identity, not a hypothesis. Written as a sum of per-modality quadratic forms, (T4) treats the modalities as conditionally independent given η . Each summand is positive semidefinite, so $\mathcal{I}_{\text{EEG+BOLD}} = \mathcal{I}_{\text{EEG}} + \mathcal{I}_{\text{BOLD}} - \mathcal{I}_{\text{prior}} \succeq \mathcal{I}_{\text{EEG}}$ holds algebraically, for every parameterization, every dataset, and every truth. “Joint is at least as informative as single-modality” therefore cannot fail, and a benchmark reporting it as a validated result reports a tautology. The non-trivial part of joint information is the EEG/BOLD cross-covariance induced by *shared process noise*, which the block-diagonal form omits; obtaining it requires whitening the stacked record rather than each modality separately. The benchmark computes both and reports their difference, which is exactly the amount by which (T4) over-counts. The discriminating comparisons are consequently native versus resampled clocks, the *size* of the second modality’s contribution to the preregistered subset, impulse versus no impulse at matched input energy, and whether any added information becomes calibrated recovery.

A continuous conduction delay is a precondition for delay identifiability. The Fisher information of τ exists only if $\partial \mu / \partial \tau$ exists. If τ is quantized to the sample grid, that derivative is identically zero and the delay carries *no* information under any design—including designs that would otherwise recover it well. A benchmark built on a quantized delay would report zero delay information for every design and mistake an implementation choice for a scientific

Program claim	Evidence that can support it	Result that weakens or falsifies it	Implementation consequence
Typed, source-native fusion is preferable to naive resampling	Better held-out likelihood, calibration, delay recovery, and intervention forecast at matched compute and parameter count, measured against a tuned resampling baseline. Adding a modality is not itself evidence: under (T4) it cannot reduce information	No reproducible gain over single-modality or carefully tuned resampling baselines; increased overconfidence or negative transfer	Retain only the provenance/type system; narrow or remove the shared latent fusion claim. <i>Narrowed on measured evidence to a claim about source-native clocks; see Section 11.5</i>
Anatomical topology improves inference	Data efficiency, out-of-distribution calibration, and causal forecast beyond dense, randomized, and distance-matched graphs	Equal-capacity controls match or exceed performance, or topology errors are absorbed by residuals	Demote anatomy from compiled constraint to weak prior for the affected scale
Multiresolution state adds information rather than decoration	Native-scale prediction, calibrated refinement, boundary agreement, and compute savings without hidden high-frequency hallucination	Fine views cannot improve supported observables, fail round-trip tests, or become overconfident outside measured tiles	Disable the scale relation or use source-specific views without gluing
Perturbation reduces non-identifiability	Rank/eigenvalue improvement in Fisher information <i>at matched input energy</i> , and prospective recovery of direction, delay, gain, dose, and state dependence	Intervention fails to distinguish posterior models, adds only field-model uncertainty, or loses its apparent advantage against an energy-matched control	Narrow the identifiable parameter set and redesign the perturbation rather than reporting a causal estimate
Individualization improves future prediction	Incremental calibrated log score and decision utility on new sessions and unseen tasks or interventions	Anatomy-only, population, or session-adapted baselines perform equivalently	Do not label the model an individual twin; retain only the supported level of adaptation

Table 1 Claims with consequences. Each central commitment has a measurement capable of making it smaller. Engineering breadth, parameter count, plausible diagrams, and in-sample fit are not substitutes for these tests. Rows 1 and 4 have been narrowed since the previous version because the measurement was run and returned a smaller result than the claim; the evidence columns record the controls that narrowing showed to be mandatory (Section 11.5).

result. (T1) and (T4) therefore require a fractional-delay implementation; here a normalized Gaussian-windowed-sinc interpolation over the history taps, so that τ is a continuous parameter with a smooth derivative. The 1 s resampled design is retained precisely because it violates this condition, and is reported as the illustration of what the violation looks like rather than as a measurement of the delay.

The benchmark compares EEG alone, fMRI alone, joint native-clock inference, joint inference after naive resampling, joint inference plus one calibrated impulse, and—required rather than optional—the same impulse at total input energy matched to the impulse-free design by scaling the background drive. Without that sixth design, an impulse comparison measures the energy budget and not the perturbation. A charitable resampling control, which discards EEG onto the 1 s grid but retains the exact fine-grained forward model, separates information lost by decimation from bias introduced by the wrong discretization. The benchmark reports rank, condition number, minimum non-prior eigenvalue, parameter-profile likelihoods, posterior correlations, interval coverage, and delay error. Prior contribution is shown separately so a full-rank posterior cannot disguise a prior-dominated likelihood.

What this system cannot be asked. The state is three regions with no spatial structure, and θ is a few coupling gains plus one conduction delay. That is an electrophysiological question, and a hemodynamic instrument is being scored on it. The strengths such an instrument actually has—spatial extent and localization, slow state, and constraining the very observation model on which electrical source inversion depends—have nowhere to appear in a three-region linear-Gaussian system. A weak fusion result obtained here is therefore a correct statement about *this design*. It is not a general finding that cross-method integration is unprofitable, and it may not be reported as one; reaching that stronger question requires extending the system, not rerunning this one.

This is a planned calculation, not a favorable result assumed in advance. The central premise is supported only if native-clock fusion or intervention increases likeli-

hood information for a preregistered parameter subset and improves calibrated recovery across held-out simulation regimes. If it does not, the compiler may still be useful as a provenance system, but the claim that cross-method integration resolves those dynamics must be narrowed. That calculation has now been performed in the linear-Gaussian case; its outcome, and the narrowing this paragraph prescribes, are reported in Section 11.5. The paragraph is left exactly as written, because it is the commitment against which the result is being judged.

The second artifact compiles the same system from a source schema and performs end-to-end recovery. It introduces known transform bias, shared session calibration error, missing windows, unequal supports, and a deliberately misspecified residual. Success requires: no leakage across simulated parent subjects; recovery intervals with nominal coverage; lower held-out log loss than single-method and naive-resampling baselines; detection of the misspecified module; and refusal of at least one invalid schema. This synthetic case is the minimum integration test for all later human experiments. It has been built and run: three of the five criteria listed here were met, interval coverage was not, the held-out log-loss comparison could not be evaluated because none of the fits it compares had converged, and a sixth subgroup-calibration criterion added by the implementation was met (Section 11.5).

0.4 Executable Evidence, Transform, and Gradient Contracts

Every dataset snapshot is compiled from a source card rather than passed to a global dataloader. A source card declares identity and hashes, participant and derivative lineage, governance, native signal and units, spatial support, point-spread function, clock, coordinate/calibration graph, observation or intervention model, missingness, validity domain, bias/variance/discrepancy ledger, permitted module and scale targets, and a fixed role: prior, likelihood, boundary target, distillation, calibration, negative control, or evaluation only. The compiler converts those declarations into dataloaders, transform paths, loss factors,

Rejected configuration	Why it is scientifically invalid	Required remedy
Unknown units, clock, physical support, frame, handedness, or transform lineage	Numerical compatibility would be mistaken for measurement compatibility	Supply a calibrated manifest with validity interval and uncertainty, or keep the source disconnected from the requested path
Prolongation without a declared restriction partner and tested coverage	A coarse observation would be converted into unsupported fine structure	Provide paired maps, round-trip and held-out landmark tests, and an out-of-support uncertainty policy; otherwise return a distribution or refuse
Global cross-scale state when overlap/cocycle residual exceeds tolerance	Conflicting local views would be hidden by one raster	Preserve separate sections, branch or marginalize the posterior, and emit an obstruction certificate identifying the failed path
Effective or causal operator estimated from passive correlation alone	Common input, observation mixing, and omitted state remain alternatives	Label the object functional/statistical, or add an identified intervention or assumption set sufficient for the causal claim
Learned residual allowed to dominate a mechanistic term silently	The named mechanism becomes unfalsifiable	Enforce a preregistered energy/gain ratio $\ \mathcal{R}_\theta\ /\ \mathcal{F}_{\text{mech}}\ \leq \rho_{\text{max}}$ on the declared validity set, report violations, or reclassify the entire module as a surrogate
Adaptive step sizes for a learned propagator without semigroup testing	Coarse and composed fine transitions need not describe the same dynamics	Measure the semigroup residual at every permitted step pair; restrict the scheduler or include the discrepancy in prediction intervals
Population, subject, and session effects without centering or shrinkage	The additive decomposition is not identified	Use centered/sum-to-zero constraints or proper hierarchical priors and test recovery in simulation
Bias assigned a point estimate without an estimator or external bound	Calibration and model discrepancy can trade off non-identifiably	Classify the term as design-estimable, externally bounded, or prior-specified sensitivity; propagate the last category as a range
Pseudo-likelihood compatibility treated as a calibrated posterior likelihood	Agreement penalties do not define a generative noise model	Report a generalized/pseudo-posterior and calibrate it empirically, or validate and promote the factor to a generative likelihood
Derived scans, sessions, relatives, stimuli, or simulator replicas crossing a parent-level holdout	Duplicated evidence produces invalid confidence and shortcut prediction	Group by immutable lineage identifiers before splitting and fail the run when parentage is unresolved
Intervention optimization outside an independently validated feasible set	A learned objective cannot guarantee device, exposure, or protocol safety	Require $u \in \mathcal{A}_{\text{safe}}$, independent safety checks, deferral, and the applicable ethics and regulatory review

Table 2 Mandatory modeling refusals. These are compiler and runtime errors, not optional governance prose. An override changes the claim carried by the resulting artifact and remains visible in its provenance.

gradient masks, and group-aware splits. An unknown field remains unknown; it is not supplied as zero or as an average brain label.

The uncertainty ledger uses three different statuses. A quantity is **design-estimable** only when the acquisition contains the replication, randomization, intervention, anchor, or hierarchy necessary to estimate it. Repeated sessions, for example, can estimate within- and between-session variance but do not by themselves identify absolute model bias. A term is **externally bounded** when a phantom, calibration target, independent instrument, or negative control supplies a defensible interval. A term is **prior-specified sensitivity** when neither is available; it is swept over a declared range and cannot be advertised as empirically estimated. Model discrepancy and calibration parameters are generally confounded without informative assumptions, a known failure mode in computer-model calibration (Kennedy and O’Hagan, 2001; Brynjarsdóttir and O’Hagan, 2014).

For a derived quantity $z = f(x, c)$, calibration uncertainty may be shared across every observation in a session. First-order propagation therefore retains cross-covariance:

$$\begin{aligned} \Sigma_z &\approx J_x \Sigma_x J_x^\top + J_c \Sigma_c J_c^\top \\ &\quad + J_x \Sigma_{xc} J_c^\top + J_c \Sigma_{cx} J_x^\top, \\ b_z &\approx J_x b_x + J_c b_c + \delta_f. \end{aligned} \quad (\text{T5})$$

The transform graph stores each edge separately. Paths are checked for units, handedness, epoch, support, round-trip residual, and calibration validity; their composed result never replaces the original edges. Figure 1 gives the frame path used by the running example.

For episode e , only authorized additive factors are active:

$$\begin{aligned} \mathcal{L} = \sum_e w_e &[\lambda_{\text{like}} \mathcal{L}_{\text{like}}^{(e)} + \lambda_{\text{boundary}} \mathcal{L}_{\text{boundary}}^{(e)} \\ &+ \lambda_{\text{distill}} \mathcal{L}_{\text{distill}}^{(e)} + \lambda_{\text{cal}} \mathcal{L}_{\text{cal}}^{(e)} \\ &+ \lambda_{\text{mech}} \mathcal{R}_{\text{mech}}^{(e)} + \lambda_{\text{scale}} \mathcal{R}_{\text{scale}}^{(e)}]. \end{aligned} \quad (\text{T6})$$

Weights depend on effective independent evidence and declared reliability, not row count. Gradient conflict is logged by source and module. A conflict may produce partial pooling, a source-specific adapter, a frozen path, a branched posterior, or rejection; consensus is not mandatory.

0.5 Running Case: Individual DLPFC TMS Targeting

The running case asks a bounded question: for one consented adult in a specified session and cognitive state, which of a small preregistered set of left-DLPFC coil poses is predicted to produce the desired short-horizon change in a measured network response while remaining inside independent exposure and protocol limits? It does not ask whether the model treats depression, nor does it infer a latent clinical construct from one signal.

1. **Declare.** The schema selects the subject’s cortical surface, uncertain tract topology, local DLPFC field, coarse distributed network, EEG and fMRI observation heads, coil geometry, neuronavigation frames, stimulation waveform, session state, and $\mathcal{A}_{\text{safe}}$.
2. **Compile.** The compiler checks source cards and transform paths, instantiates a fine DLPFC tile plus coarse surrounding regions, registers restriction/prolongation maps, and emits native EEG and fMRI clocks. A pose expressed only as “5 cm anterior” is rejected.
3. **Calibrate.** MRI constrains geometry; digitized electrodes

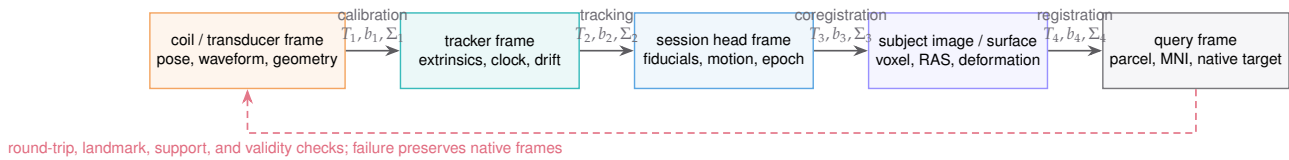


Figure 1 Coordinate and calibration frame graph. Device commands, tracked pose, head coordinates, subject images, and atlas or query coordinates remain distinct nodes connected by versioned transforms. Each edge retains calibration observations, epoch, systematic offset, random covariance, and model discrepancy. Composition is query-specific; raw transforms are never multiplied away, and a failed cycle prevents a global coordinate claim.

and head model constrain the EEG lead field; tracker-to-coil and head-to-MRI observations constrain pose; motor-threshold or other approved calibration informs device gain without being equated to DLPFC target engagement.

4. **Infer.** Baseline recordings produce a posterior over neural state, coupling, observation nuisance, tissue parameters, and model class. Session calibration errors remain correlated rather than being averaged away. Models that fit passive data but disagree about propagation are retained.
5. **Counterfactually compare.** Each admissible pose is passed through the electromagnetic field solver and candidate neural response operators. The system predicts a distribution over rapid EEG propagation, later hemodynamic consequences, burden, and explicit harm endpoints. Field accuracy, target engagement, network effect, and clinical utility remain separate quantities.
6. **Decide or defer.** The optimizer acts only within $\mathcal{A}_{\text{safe}}$. If model disagreement or transform uncertainty dominates the estimated benefit difference, the output is another calibration measurement, a reversible probe, or no recommendation.
7. **Test.** A prospective, blinded or otherwise causally identified comparison evaluates predicted direction, latency, spatial propagation, dose response, state dependence, calibration, and decision regret. The result updates only the parameters and claims authorized by the protocol.

Prospective human TMS or tFUS work requires ethics, safety, consent, and device review appropriate to the actual protocol. In the United States the sponsor and IRB must document significant-risk versus nonsignificant-risk status; a significant-risk device investigation requires FDA and IRB approval, while an NSR study follows the abbreviated IDE requirements after IRB concurrence ([U.S. Food and Drug Administration, 2006](#)). IRB approval alone is therefore not treated as a universal regulatory description.

0.6 Agent-Facing Build Order and Definition of Done

Implementation proceeds through claim-bearing vertical slices rather than a simultaneous attempt at every brain system.

1. **Schema kernel.** Implement typed regions, ports, operators, clocks, supports, frames, authority policies, uncertainty status, source roles, and immutable lineage. Definition of done: valid examples compile; every refusal in Table 2 has a failing fixture.
2. **Linear identifiability laboratory.** Implement Equations (T1)–(T4), analytic/numerical Fisher information, simulation, posterior recovery, and five design comparisons. Definition of done: deterministic seeds, parameter sweeps, coverage tests, and a machine-readable claim report.
3. **Transform and clock runtime.** Implement frame-graph composition, cross-covariance, multirate events, validity intervals, and obstruction certificates. Definition of done:

unit/handedness/clock errors are caught; round trips and shared calibration error are tested.

4. **End-to-end synthetic slice.** Compile three regions, EEG/fMRI reads, one impulse write, missing windows, and model discrepancy. Definition of done: the preregistered baselines and acceptance criteria in Section 0.3 run from one manifest.
5. **Empirical subsystem.** Add one anatomically bounded human subsystem with public data, source-native likelihoods, and held-out cross-modal prediction. Definition of done: provenance, leakage, bias status, and negative-transfer ablations ship with the result.
6. **Prospective perturbation pilot.** Only after solver, field, calibration, and causal recovery gates pass, instantiate the running TMS or tFUS case under an approved protocol. Definition of done: preregistered prospective predictions and an explicit no-benefit/no-identification outcome path.

Every implementation agent receives the thesis version, schema version, target claim, allowed source cards, acceptance tests, and non-goals. A merged component must add tests, units and frame declarations, uncertainty behavior, provenance fields, baseline comparisons, and a statement of what empirical finding would disable it. Code that only expands the operator registry is not scientific progress until it passes a claim-bearing gate.

1. Introduction

The central difficulty of whole-brain modeling is not scale alone but *incommensurability*. Structural MRI constrains millimetric anatomy; diffusion MRI provides uncertain macroscopic tract evidence; fMRI samples a spatially blurred vascular response over seconds; EEG and MEG sample rapid electromagnetic mixtures whose cortical source inversion is non-unique; TMS introduces an orientation-dependent electric field; and transcranial focused ultrasound (tFUS) introduces an acoustically transformed focal exposure. These measurements and perturbations refer to one nervous system, yet they do not inhabit one raster, one clock, one coordinate frame, or one ontology. Multimodal datasets already require fiducial digitization, MRI coregistration, distinct forward models, and asymmetric fusion even for a single controlled task ([Henson et al., 2019](#)). For individualized neurotechnology, the problem is harder: the model must predict how a millisecond electrical trajectory, a second-scale hemodynamic trajectory, a device pose, a focal physical field, and a later behavioral outcome constrain one another without erasing their native supports or uncertainties.

This is the direct utility proposed for SuperCognition Whole-Brain Dynamics (SC-WBD). MRI can constrain individualized geometry and slow observation physics; EEG can update fast latent dynamics; TMS and tFUS can act as causal writes with their own physical point-spread functions; and all four can be joined through an explicit latent process and calibrated transforms rather than by resampling them into

fictitious equivalence. A model may therefore answer cross-method questions: which rapid state was likely present during an fMRI response, which network consequences should follow a TMS field placed at a recorded coil pose, whether a tFUS focus overlaps an inferred causal bottleneck after skull transmission, and which additional measurement would most reduce decision uncertainty. These are conditional predictions, not claims that one modality recovers the hidden ground truth of another.

The first measured version of this argument is narrower than the argument itself, and the narrower version is the one this paper now asserts. In the linear-Gaussian laboratory of Section 0.3, what survives measurement is the handling of *clocks*, not the addition of *modalities*. Resampling a millisecond record onto a one-second lattice destroys delay information that remains present in the data, whereas a native-clock estimator retains it. The second modality's own contribution, in that small system, is about one percent of the electrical channel's information about the coupling gains and—because a millisecond delay does not survive a multi-second hemodynamic convolution—essentially none about the conduction delay; and a calibrated impulse's apparent advantage is accounted for by its input energy (Section 11.5). Adding a modality is not automatically beneficial, and nothing below assumes that it is. Nor does that measurement license the converse: the benchmark asks a question shaped to one of its two modalities, and says nothing about the strengths the other one has no room to exhibit there.

Whole-brain network models have shown that empirical anatomy can constrain large-scale dynamics and generate subject-specific observables (Sanz Leon et al., 2013; Sanz Leon et al., 2015; Jirsa et al., 2017). Work on integrative and scale-sensitive modeling likewise argues that structure and dynamics interact across micro-, meso-, and macroscales (Venkadesh and Van Horn, 2021; Kobeleva et al., 2021). Yet fidelity and identifiability remain in tension. One neural mass per parcel is tractable but can erase within-region geometry, laminar routing, frequency structure, and population heterogeneity. A uniformly microscopic reconstruction preserves detail that present non-invasive measurements cannot identify and makes posterior inference prohibitive. SC-WBD consequently does not assign the brain one spatial or temporal resolution. Resolution is a property of each state, source, operator, intervention, and query.

The basic computational object is an operator graph over heterogeneous structured state spaces. A cortical area may contain a surface field, laminar and cell-class channels, spectral variables, spike events, sparse memories, hemodynamic compartments, and uncertainty fields rather than one activation scalar. A self-connection or cross-region connection is not one weight: it is a typed operator with an admissible function family, spatial footprint, coordinate convention, temporal support, delay, plasticity, evidence class, and bias-variance ledger. Local convolutional, mesh, and neural-field topologies coexist with sparse long-range projections. At each self-edge or cross-edge the implementation can be a biophysical flow, delayed kernel, state-space channel, spectral transfer function, attention map, point-process transform, learned surrogate, or calibrated composition. Architectural generality is not neuroscientific evidence; a function family earns utility only when its specific commitments improve held-out, cross-modal, or perturbational predictions.

Machine learning supplies neural differential equations, graph neural fields, state-space models, sparse attention, diffusion-based inference, and learned simulators (Aqil et al., 2021; Kashyap et al., 2023). Their capacity is valuable, but ordinary prediction can be achieved by a latent organization that conflicts with adult anatomy or fails under

intervention. Conversely, a visually plausible anatomical simulation may be unidentifiable and unfalsifiable. SC-WBD uses anatomy as a probabilistic interaction grammar, physiology as operator priors and observation physics, and learning as the mechanism that fits local computations and effective routing. Its compiler makes each commitment pointable so it can be ablated or replaced.

The architecture is human-brain specific at its core. Atlas families, cortical geometry, laminar evidence, tract priors, receptor maps, body interfaces, behavioral paradigms, and intervention operators are organized around the adult human nervous system. Comparative animal data may constrain conserved mechanisms or expose failures of transfer, but the initial ontology is not a generic mammalian brain. SC-WBD also does not simulate morphogenesis. Genetic, transcriptomic, and developmental evidence is used only where it improves priors over a mature regional phenotype, projection grammar, receptor topography, intrinsic timescale, or admissible connectivity. Within the adult model, fast state evolves over milliseconds to seconds, effective coupling and synaptic efficacy learn over experience, and local meso- or microscale structure may revise conservatively over weeks to months.

Nor must development wait for a synchronized whole-brain-whole-body dataset. Public, private or partnered, prospectively acquired, and simulated episodes can supervise only the subgraph, modality port, body interface, and native scale they actually observe. Eye movements alone constrain an oculomotor slice; eye movements paired with visual content additionally constrain retinotopic and salience dynamics; embodied simulations can constrain motor, proprioceptive, tactile, auditory, visual, nociceptive, and interoceptive interfaces; physiological recordings can constrain autonomic and digestive control. Missing regions are latent, frozen, sampled at their boundaries, or marginalized—not assigned invented labels. Shared typed ports permit regional models learned from different studies to become mutually constraining when assembled.

The objective is therefore an individualized posterior over a **multiresolution adult brain process**: heterogeneous dynamical systems coupled by explicit operators, observed through source-specific models, registered through uncertain coordinate chains, and perturbed through causal intervention models. Every object retains estimated systematic bias, stochastic variance, provenance, and a validity domain. Near-term horizons are cross-modal forecasting, source-aware EEG control, and offline comparison of stimulation hypotheses; intermediate horizons require prospective prediction of individual perturbational response; longer-term person models require grounded world, body, self, affective, memory, social, and language dynamics. The horizons are separated by validation gates rather than compressed into one promise.

2. Compiled Structured-State Architecture

2.1 Regions Are Structured State Spaces

For subject p , SC-WBD defines an operator-valued simulator graph

$$\mathcal{G}_p = \left(\mathcal{V}, \mathcal{E}, \{\mathcal{X}_i\}_{i \in \mathcal{V}}, \{\mathcal{K}_{ij}\}_{(i,j) \in \mathcal{E}}, \boldsymbol{\tau}, \boldsymbol{\Pi}_p \right),$$

where $\mathcal{V}$ indexes modeled systems, $\mathcal{X}_i$ is the state space owned by region i , $\mathcal{K}_{ij}$ is the directed operator from i to j , $\boldsymbol{\tau}$ contains conduction and temporal-support priors, and $\boldsymbol{\Pi}_p$ contains individual parameters and uncertainty. A cortical region may, for example, maintain

$$X_i(t) = (X_i^{\text{sheet}}, X_i^{\text{layer}}, X_i^{\text{population}}, X_i^{\text{frequency}}, \\ X_i^{\text{memory}}, X_i^{\text{metabolic}}, X_i^{\text{uncertainty}}) \in \mathcal{X}_i.$$

The components need not have equal shape or even be ordinary dense tensors. They may be fields on a cortical mesh, arrays indexed by depth and cell class, graphs of local populations, point processes of spike events, sparse distributed codes, sets of particles, probability distributions, or low-dimensional sufficient statistics. A module can expose a coarse boundary state while retaining a much larger private state. Conversely, a scientific query may refine one boundary into a richer interface when the necessary data and computation are available.

The regional update is written explicitly as

$$X_i(t + \Delta t_i) = \Phi_i^{\Delta t_i} \left(X_i(t), U_i(t), C(t); \right. \\ \left. \{ \mathcal{K}_{ji}[X_j]_{t-\tau_{ji}:t} \}_{j \rightarrow i}, \theta_i \right). \quad (1)$$

Here Φ_i may be a numerical flow, stochastic transition, event-driven update, learned neural operator, or composition of these; U_i is an intervention or bodily input; and C is task and environmental context. Equation (1) is deliberately representation agnostic, but not semantically agnostic. Every state component and operator declares units, spatial support, temporal support, uncertainty, and provenance. Figure 2 illustrates the distinction between a structured region and an operator-valued connection.

Self-dynamics and cross-region communication are dispatched from the same typed operator registry:

$$\mathcal{K}_{ij} \in \{ \text{flow}_{\text{ODE/PDE}}, \text{field kernel}, \\ \text{convolution, delayed SSM, spectral transfer}, \\ \text{attention, point process, surrogate, composition} \}. \quad (2)$$

For $i = j$, the operator may express local recurrence, reaction–diffusion, homeostatic control, or intrinsic memory; for $i \neq j$, it maps a declared source boundary into a compatible destination port. Each instance records its input and output type, coordinate frame, clock, differentiability, solver, free and fixed parameters, mechanistic status, evidence base, and bias–variance decomposition. Allowing an arbitrary transfer family at each self- or cross-connection is a *model-class capability*. It becomes a neuroscientific contribution only when alternative operators are compared on measurements or perturbations that discriminate their dynamics.

2.2 The Connectome as a Topology Prior, Not a Scalar Matrix

At one useful abstraction, an anatomical connectivity relation becomes a source-by-destination block topology over latent state. Block-diagonal structure preserves dense local interaction within a field or network, whereas selected off-diagonal blocks permit specific long-range messages. The blocks may differ in shape, update frequency, internal dimensionality, operator family, and loss role. This is more informative than exposing every latent element to every other element and asking data alone to rediscover the interaction graph. It also makes the model pointable: each allocated block has a declared interpretation, and every permitted communication path can be inspected, disabled, replaced, or subjected to model comparison.

The analogy should not be literalized. A tractography streamline count is not an attention weight, a functional-correlation matrix is not a connectome, and the absence of

an observed tract is not proof of conditional independence. SC-WBD therefore assigns pathways to three compilation classes:

1. **hard-supported edges** represent repeatedly established pathways whose removal would contradict the chosen anatomical model;
2. **soft-supported edges** represent uncertain, indirect, subject-variable, or method-dependent evidence and carry shrinkage priors;
3. **proposed edges** are admitted only within an explicit model comparison and pay a complexity, distance, and provenance penalty.

A zero mask creates exact independence across that direct edge inside the compiled model. It does not establish biological independence because communication may occur through another path, a shared input, volume conduction, neuromodulation, vascular coupling, or an omitted population. Likewise, a permitted edge does not imply that it is causally active in every state. The effective operator may be gated by oscillatory phase, thalamic or basal-ganglia control, task context, neuromodulators, and local excitability.

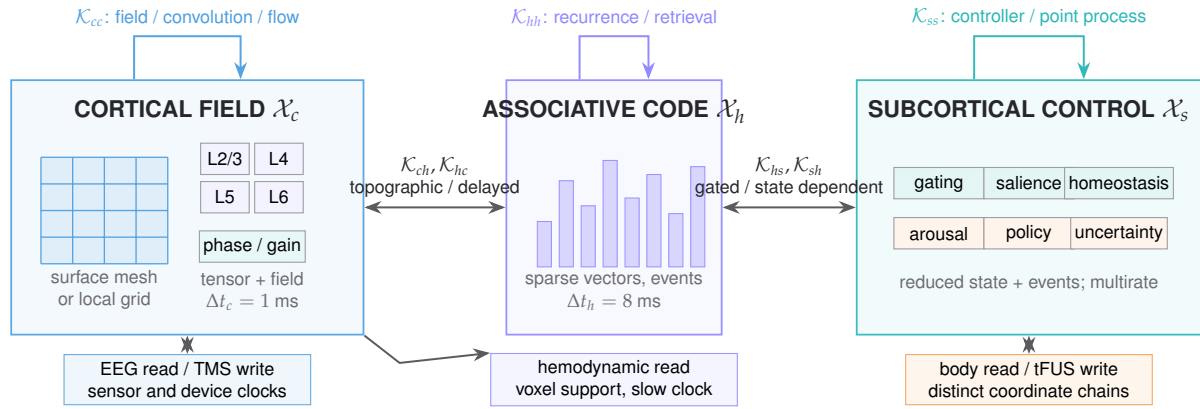
The edge operator $\mathcal{K}_{ij}$ can instantiate block-sparse attention, a delayed neural-field kernel, a graph convolution, a low-rank linear map, a state-space channel, a conductance-informed projection, a stochastic point-process transform, or a mixture selected by model evidence. The phrase every cell has an operator type is therefore a useful compiler intuition, provided that a cell is understood as a typed map between boundary spaces rather than as a single learned coefficient.

2.3 The Brain Schema as Memory Map and Calling Convention

SC-WBD uses a declarative schema to describe anatomical entities, nested regions, state fields, ports, pathways, clocks, observation roles, interventions, and confidence. A compiler then emits:

- packed latent-memory regions and stable offsets for every state field;
- a resolution contract for each field: native supports, authority policy, calibrated cross-scale maps, gradient routes, and refinement triggers;
- local grid, mesh, or convolutional neighborhoods within spatial regions;
- block-sparse cross-region attention masks and typed operator dispatch;
- temporal frames, delays, update frequencies, persistence rules, and multirate synchronization points;
- loss roles, priors, conservation or homeostatic penalties, and separate bias, variance, and model-discrepancy heads;
- observation adapters for instruments and intervention adapters for causal input;
- a serialized state schema through which compatible models can exchange state without re-tokenizing it.

This resembles a memory map, application binary interface, and calling convention for latent dynamics. Two implementations that share the same schema can exchange documented boundary state even when one uses a neural field and the other a learned surrogate internally. The symbolic component is the engineering declaration itself: region identity, edge existence, type, direction, and clock. Neural computation remains inside the structured states and operators. There is no mandatory discrete-symbol interface between symbolic macrostructure and neural microstructure because the macrostructure compiles directly into the permissible computation. A prior software prototype informed the declaration-to-layout intuition (Valdez, 2026a);



Every state block and self/cross operator declares shape, support, clock, units, coordinates, evidence, plasticity, mechanistic status, bias, variance, and provenance.

Figure 2 Heterogeneous structured regions and operator-valued connections. Allocated regions need not have the same shape, dimensionality, clock, or internal topology. Each self-connection and cross-region connection selects its own transfer family and maps a documented boundary representation into a compatible destination port. Observation and intervention ports attach at their physical supports and clocks. The drawing shows three regions and three empirical ports for clarity; the compiled graph may contain arbitrarily sized and nested blocks at many resolutions.

this is a provenance citation, not neuroscientific evidence, and no claim in SC-WBD depends on that prototype's terminology or reported performance.

The architecture in Figure 3 separates evidence and priors, compiled structured dynamics, the assimilated latent process, and empirical interfaces. This separation makes it possible to disable expensive physics without altering the conceptual model. A run may evolve only latent neural states; add electrophysiological source and volume-conduction models; add hemodynamics; generate tractometry-related predictions; or locally refine a cortical patch for a perturbational question.

2.4 Latent Process, Observation, and Intervention

The latent process must not be defined by any one instrument, and instruments need not be coerced into a common resolution. For modality m ,

$$Y_m(t) = \mathcal{O}_m(X_{0:t}, A_p, B_p; \psi_m) + \epsilon_m(t),$$

where $\mathcal{O}_m$ is an observation operator, A_p is subject anatomy, B_p is body and behavioral state, and ψ_m includes device and session parameters, native support, and point-spread uncertainty. EEG and MEG heads project current distributions through subject-specific forward models. fMRI and fNIRS heads introduce neurovascular and metabolic states before sampling. Diffusion MRI and tractometry constrain structural priors or predict slow tissue-associated observables; they are not treated as rapidly changing neural activation. Behavior and report are generated through perception, policy, motor execution, language, and reporting bias. Native observations and their residuals are retained even when the inference system proposes a common latent explanation.

An intervention is separately represented as a controlled input over its physical exposure interval:

$$\begin{aligned} dX(t) &= \mathcal{F}(X, t) dt \\ &+ \mathcal{G}_k(X, t; A_p, C, \omega_k) u_k(t) dt \\ &+ Q^{1/2} dW_t, \quad t \in [t_0, t_1]. \end{aligned}$$

Device geometry, waveform or burst sequence, cumulative thermal history, tissue coupling, and mechanistic uncertainty remain distinct. An instantaneous jump $X(t_0^+) = \mathcal{I}_k(X(t_0^-), \int u_k dt, \dots)$ is admitted only as an explicitly tested impulse limit, for example for a sufficiently brief TMS pulse.

This prevents an induced electric field, acoustic pressure, or presented sentence from being equated with a neural effect. It also permits causal calibration: two latent models that fit resting data equally well may predict different responses to the same controlled perturbation.

2.5 Structural, Effective, and Functional Connectivity

Three matrices that are often conflated occupy different layers:

$$S_{ij} = \Pr(\text{anatomical pathway } i \rightarrow j \mid \text{evidence}),$$

$$E_{ij}(t, C) = \text{state- and context-dependent effective operator},$$

$$F_{ij}^{(m)} = \text{stat}(Y_i^{(m)}, Y_j^{(m)}).$$

The first is a structural prior, the second is part of the generative dynamics, and the third is a modality- and estimator-dependent statistic. Functional connectivity can change without a tract changing; effective coupling can reverse or gate information flow over a fixed tract; an apparent functional edge can arise from a common driver. The compiler records these distinctions so that an anatomical mask does not quietly become a causal claim.

2.6 A Resolution Poset Rather Than a Unified Raster

Resolution is a coordinate of model state, evidence, and intervention. Let $\mathfrak{R} = (\mathcal{R}, \leq)$ be a partially ordered set of spatial, temporal, anatomical, and representational supports. Its elements need not be dyadic image sizes or even grids: they may be surface meshes, parcels, laminar depths, voxel fields, tract endpoint distributions, spectral bands, event windows, behavioral episodes, or sparse local tiles. An empirical or simulated source s contributes a typed partial view

$$Q_s = (D_s, r_s, K_s, \Delta t_s, u_s, \Sigma_s, \text{prov}_s),$$

where D_s is its domain, $r_s \in \mathfrak{R}$ is native resolution, K_s is its point-spread or support kernel, u_s denotes units, Σ_s denotes uncertainty, and prov_s records provenance. Every read and write request carries the same metadata.

Mappings between views are calibrated operators, not automatic resampling. For compatible scales $r_a \leq r_b$, a restriction $\mathcal{R}_{a \rightarrow b}$ may be available; a prolongation $\mathcal{P}_{b \rightarrow a}$ may produce a distribution rather than a unique fine state. Some

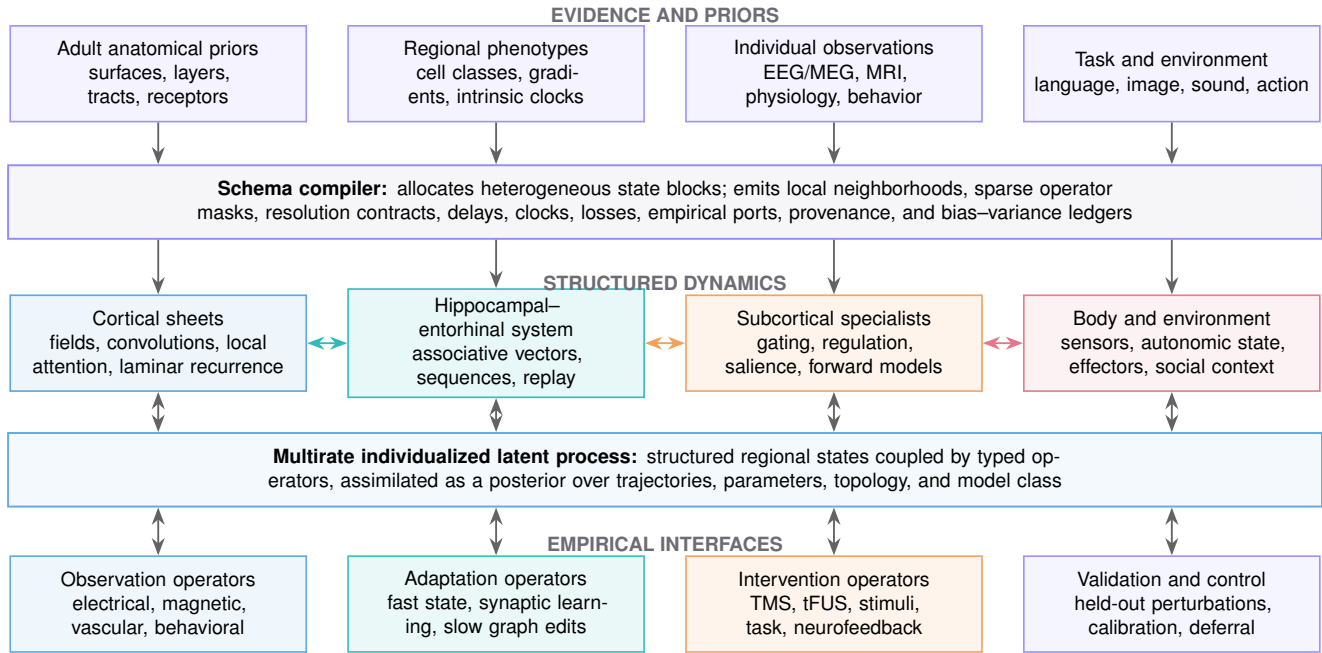


Figure 3 Compiled whole-brain modeling framework. Anatomical and phenotypic evidence defines a probabilistic adult-brain schema; individual observations and task context constrain its subject-specific state. A compiler maps the schema into heterogeneous memory regions, local cortical neighborhoods, sparse long-range operators, update clocks, and empirical interfaces. Runtime modules need not share a representation or resolution. Their contract is an explicit state type, operator family, spatial and temporal support, uncertainty semantics, and a validated boundary map.

pairs have no defensible map. Multiple sources may therefore disagree on an overlap or fail to admit a single globally consistent field. The runtime preserves those partial views and their residuals rather than forcing every input into a unified truth tensor.

The design can be made operational rather than merely “sheaf-like.” Let the site $\mathcal{C}$ contain declared domain–support objects and admissible covers; let $\mathcal{F}(U)$ be the set or distribution of states admissible on U ; and let every inclusion $V \hookrightarrow U$ own a versioned restriction ρ_V^U . Local sections $x_a \in \mathcal{F}(U_a)$ are compatible only when their overlap residual is inside a declared uncertainty tolerance. Alternative paths must also approximately commute:

$$\Omega_{abc}(x) = \left\| \rho_{U_c}^{U_a} x - \rho_{U_c}^{U_b} \rho_{U_b}^{U_a} x \right\|_{W_c} \leq \tau_{abc}.$$

If overlap and cocycle tests pass, the runtime may construct a global section and records the gluing residual. If they fail, the nonzero residual is an *obstruction certificate*: the runtime preserves separate views, branches or marginalizes the posterior, or refuses the global materialization. Thus a shared latent process is a generative hypothesis earned through cross-source prediction and commuting transforms, not imposed by upsampling.

2.7 A Bias–Variance Ledger for Every Computational Object

Uncertainty is not represented by one confidence score. Every source, latent field, parameter, edge, solver, cross-scale map, observation operator, and intervention prediction carries separate metadata for systematic discrepancy and stochastic variability. For source s ,

$$Y_s = \mathcal{O}_s(X; D_s, r_s) + b_s(D_s, r_s, C) + \varepsilon_s, \\ \mathbb{E}[\varepsilon_s] = 0, \quad \text{Cov}(\varepsilon_s) = \Sigma_s.$$

The bias field b_s can include device calibration error, atlas mismatch, aggregation or partial-volume bias, dataset and sampling bias, omitted physiology, task confounding, and stable model discrepancy. It is itself uncertain and is

assigned a prior or empirical interval; labeling a term bias does not mean its true value is known. The variance ledger separates, where identifiable, measurement noise, within-session stochasticity, between-session biological variation, parameter-posterior variance, model-class disagreement, and numerical or Monte Carlo error.

Cross-scale maps have the same contract:

$$\hat{X}^{(r_b)} = \mathcal{T}_{r_a \rightarrow r_b} \left[X^{(r_a)} \right] + b_{a \rightarrow b} + \varepsilon_{a \rightarrow b}.$$

A coarse summary can therefore have low variance but high aggregation bias; a fine reconstruction can have low estimated bias inside a measured tile and very high variance outside it. Losses are not automatically trusted merely because their variance is small. Each systematic term receives one of three statuses: *design-estimable* when replication, randomization, intervention, anchors, or hierarchy identify it; *externally bounded* when phantoms, calibration targets, negative controls, or independent benchmarks constrain it; or *prior-specified sensitivity* when neither is available. Repeated sessions primarily estimate variance components; nested cross-validation estimates generalization error; modality disagreement bounds relative discrepancy only under an explicit error model. None alone supplies an absolute bias estimator. Calibration parameters and model discrepancy are often jointly non-identifiable without informative assumptions (Kennedy and O’Hagan, 2001; Brynjarsdóttir and O’Hagan, 2014).

Each read returns a prediction, variance decomposition, estimated bias range, model-discrepancy flag, provenance, and validity domain. Each write returns the same ledger for delivered exposure and predicted effect. Bias and variance are propagated analytically where possible and by ensembles, posterior sampling, or sensitivity analysis otherwise. Population and subgroup audits are retained separately: a globally calibrated model may still be biased for a demographic, anatomical, clinical, device, or task subgroup that was poorly represented during development.

2.8 Coordinate Frames, Calibration Points, and Pose Uncertainty

A location is not a three-number property of a brain. It is a point or pose in a named reference frame, measured through a calibration chain at a particular session and time. SC-WBD therefore represents every sample and impulse as

$$\mathcal{C} = (q, \text{frame}, \text{fiducials}, t, \text{units}, K, b_{\mathcal{C}}, \Sigma_{\mathcal{C}}, \text{prov}),$$

where q may include position and orientation, K is the physical support or point-spread function, and the remaining terms state calibration, bias, covariance, and provenance. Five millimeters from left DLPFC is underspecified unless DLPFC is defined on a particular individual surface or atlas and the surface-to-device transform is known. Likewise, pitch and yaw are meaningful only relative to a stated coil, tracker, head, or nasion–inion–preauricular frame.

For session s of person p , a representative device-to-atlas chain is

$$\begin{aligned} T_{p,s}^{\text{atlas} \leftarrow \text{device}} &= T_p^{\text{atlas} \leftarrow \text{image}} T_{p,s}^{\text{image} \leftarrow \text{head}} \\ &\cdot T_s^{\text{head} \leftarrow \text{tracker}} T_s^{\text{tracker} \leftarrow \text{device}}, \quad T \in \text{SE}(3). \end{aligned} \quad (3)$$

The transforms are stored separately rather than multiplied and forgotten. Their uncertainties include fiducial localization, scalp digitization, head–MRI coregistration, atlas normalization, optical-tracker calibration, instrument geometry, head motion, operator drift, and time synchronization. For small perturbations, pose covariance is propagated through the chain with the relevant Jacobians or adjoint maps; systematic offsets are propagated as a separate twist bias. Non-rigid image-to-atlas warps retain local Jacobians and registration residuals rather than being disguised as rigid pose error.

An EEG electrode is first a contact in a cap or digitizer frame. It acquires a subject-MRI location only after digitization and coregistration, and it does not become a cortical MNI source coordinate without an explicit head model, inverse operator, and atlas warp. TMS similarly requires the full coil pose, its orientation relative to the local cortical normal and tangents, and the electric-field solution—not a scalp label alone (Caulfield et al., 2022; Uehara et al., 2024). A tFUS target additionally traverses transducer, tracker, head, MRI or CT, and acoustic-simulation frames; skull properties, coupling, steering, and focal displacement contribute uncertainty (Phipps et al., 2024). Repeated calibration points and raw transform records are therefore first-class training data.

3. Adult Anatomical Priors and Plasticity

3.1 Compiling Mature Organization

Adult priors are assembled from cortical surfaces and parcellations, cytoarchitectonic and transcriptomic gradients, receptor and myelin maps, laminar projection evidence, tractography, comparative tracing, and population variation. Multimodal parcellation demonstrates that cortical boundaries can be informed jointly by architecture, function, connectivity, and topography (Glasser et al., 2016); multi-layer connectome analyses further warn that a partition appropriate to one relation may not be sufficient for another (Albers et al., 2022). SC-WBD therefore supports nested and overlapping region definitions rather than treating one atlas as ground truth.

Direction and laminar termination are especially uncertain in the human macroscale connectome. Tracer studies in

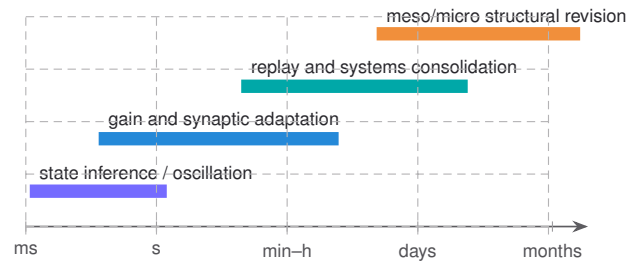


Figure 4 Nested adaptive timescales. The architecture distinguishes fast latent dynamics, intermediate gain and synaptic changes, consolidation, and conservative structural revision. Ranges overlap and depend on subsystem and task; the bars are scheduler classes rather than universal biological constants.

non-human primates provide weighted, directed interareal evidence (Markov et al., 2014), while human diffusion MRI supplies subject-specific but generally undirected and error-prone tract estimates. Mouse mesoscale resources demonstrate the value of systematic whole-brain tracing but do not license a direct species transfer (Oh et al., 2014). Each compiled edge consequently retains species, method, directionality, resolution, and confidence metadata.

Genetic and morphogenetic constraints enter only insofar as they improve the adult prior. Genetic programs contribute to circuit assembly, while experience-dependent activity refines connectivity (Luo, 2021); common genetic variation is also associated with measurable structural-connectome variation (Wainberg et al., 2024). These findings motivate hierarchical priors, not genomic determinism. The model does not infer a person’s mental organization from genotype, nor does it require simulating embryonic tissue to initialize adult cortex.

3.2 Stable Scaffold, Adaptive Routing

The macro-connectome is treated as a slowly changing scaffold rather than an immutable graph. Most learning can be expressed through changes in synaptic efficacy, short-term plasticity, inhibition, gain, synchronization, terminal arborization, and routing over existing pathways. Adult axonal boutons and local connections nevertheless remain structurally dynamic (Stettler et al., 2006). SC-WBD therefore distinguishes:

$$\theta(t) = (\theta^{\text{state}}, \theta^{\text{gain}}, \theta^{\text{synapse}}, \theta^{\text{structure}})$$

with progressively stronger priors and slower update clocks. A candidate structural edit is proposed only when a persistent residual cannot be explained by state, measurement, or effective-coupling uncertainty. It is then evaluated during replay against previously acquired competencies, energetic costs, anatomical plausibility, and stability. This is closer to constrained graph rewriting than to ordinary gradient descent over a fully dense matrix.

The update classes in Figure 4 are intentionally broad. Biological timescales overlap, and learning can rapidly alter effective connectivity without creating a new mesoscale tract. The model therefore reports which parameter class changed rather than describing every improvement as rewiring.

4. Hybrid Cortical Dynamics

4.1 Local Cortical Fields and Sparse Long-Range Operators

The cortical sheet is simultaneously local and nonlocal. Nearby populations interact through horizontal, vertical, and recurrent circuitry embedded in a continuous folded

Method	Native model port	Spatial-temporal support	Reference and calibration chain	Dominant bias-variance terms
Structural and diffusion MRI	Anatomy, tissue, surfaces, uncertain tract prior	Voxel or surface support; static or longitudinal	Scanner coordinates → subject image → surface or atlas	Gradient distortion, motion, segmentation, susceptibility, tractography false positives/negatives, atlas warp
fMRI	Neural-to-hemodynamic observation operator	Voxel point-spread plus vascular support; repetition-time and second-scale response	Scanner → functional image → anatomical image → subject/atlas	Motion, physiology, slice timing, susceptibility, neurovascular variability, partial volume, session drift
EEG/MEG	Electromagnetic forward and inverse operators	Sensor mixtures at millisecond sampling; spatial resolution is model dependent	Sensor or digitizer → fiducial-defined head → MRI; optional source → atlas	Contact/helmet location, impedance, conductivity, reference, artifacts, coregistration, source non-identifiability
TMS	Coil pose → electric field → state-dependent neural-response operator	Pulse waveform and distributed cortical field; effect depends on orientation and state	Coil → tracker → head/fiducials → MRI/surface/target	Pose and angle error, motion, coil geometry, tissue conductivity, field-model error, excitability and dose variability
tFUS	Transducer drive → in situ acoustic field → candidate tissue operators	Three-dimensional focus, pulse sequence, thermal and mechanical exposure histories	Transducer → tracker/head → MRI/CT → acoustic grid/target	Skull and coupling model, registration, steering, focal shift, cavitation/thermal estimates, unresolved neuromodulatory mechanism

Table 3 Source-native interfaces for cross-method whole-brain inference. The architecture relates methods through explicit observation or intervention operators rather than assigning them one nominal resolution. Spatial support, clock, coordinate chain, calibration points, bias, and variance remain part of every read or write.

surface; selected long-range projections connect distant areas. SC-WBD represents this with two complementary topologies.

Within a cortical patch, activity may evolve on a surface mesh or local grid using convolution, graph diffusion, neural-field kernels, local attention, reaction-diffusion terms, or laminar recurrent circuits. Locality preserves retinotopy, tonotopy, somatotopy, cortical distance, traveling waves, and spatially organized receptive fields. Graph neural-field formulations already show how stochastic integro-differential dynamics can be placed on empirical connectomes rather than restricted to Euclidean grids (Aqil et al., 2021).

Across patches, the compiler instantiates sparse directed operators informed by the connectome. A destination block may attend to a boundary summary of the source, receive a delayed field projection, or exchange a low-rank state-space message. A long-range operator need not expose every source cell to every destination cell. Its spatial footprint can encode topographic mapping, laminar origin and termination, receptor-specific gain, or an uncertainty distribution over these properties.

This hybrid produces a useful factorization:

$$\mathcal{F}(X) = \mathcal{F}_{\text{local}}(X; \mathcal{M}) + \mathcal{F}_{\text{long}}(X; \mathcal{G}) + \mathcal{R}_{\theta}(X, C),$$

where $\mathcal{M}$ is the local mesh and $\mathcal{G}$ is the long-range graph. The learned residual is regularized for magnitude, locality, intervention behavior, and out-of-distribution uncertainty so that it cannot silently replace the mechanistic terms.

4.2 Scale-Indexed State, Adaptive Refinement, and Boundary Sufficiency

A uniform resolution is neither necessary nor defensible. During a resting whole-brain run, an association region might be represented by tens of latent field modes. During TMS simulation, the stimulated gyrus and its immediate targets might be refined into laminar fields with subregional geometry. During a hippocampal memory experiment, an allocortical module might be expanded into sparse populations while the rest of the brain remains coarse.

Refinement is valid only if the boundary representation is sufficient for the surrounding model. Where the resolution

poset admits the relationship, SC-WBD learns and validates paired restriction and prolongation operators:

$$X_i^{\text{fine}} \xrightleftharpoons[\mathcal{R}_i]{\mathcal{P}_i} X_i^{\text{coarse}},$$

where $\mathcal{R}_i$ compresses fine state into the coarse sufficient statistics consumed by neighbors and $\mathcal{P}_i$ expands incoming coarse messages into a conditional distribution over fine states. Cross-resolution consistency is assessed by whether a refined module preserves observable and perturbational predictions at its boundary, not merely whether its internal trajectory looks plausible.

Resolution may be simultaneous rather than substitutive. The complete private regional state remains $X_i \in \mathcal{X}_i$. Its scale/source collection $\mathcal{S}_i$ consists of typed views $V_{i,s}^{(r)}(X_i)$, or of independently parameterized consensus states explicitly coupled to that private state. A single cortical patch can therefore maintain any finite set

$$\mathcal{S}_i = \left\{ V_{i,s}^{(r)}(X_i) \in \mathcal{X}_{i,s}^{(r)} : r \in \mathfrak{R}_i, s \in \mathcal{A}_{i,r} \right\},$$

where s indexes sources or model fields available at resolution r . There is no requirement that scales be dyadic, isotropic, nested, or shared by all sources. Geometry-preserving pooling, graph coarsening, wavelet or Laplacian transforms, learned operators, and probabilistic forward models are registered only where their validity has been tested.

Each state field declares one of three authority policies. In a **fine-authoritative** field, an $N \times M$ or mesh-level state owns the degrees of freedom and coarser views are differentiable materializations. A loss at $N/a \times M/b$ descends through the adjoint of the calibrated projection. In a **consensus multilevel** field, several scales own degrees of freedom simultaneously and exert bidirectional gradient pressure through optional compatibility potentials. In a **coarse-authoritative field with sparse refinement**, the global state lives mainly at a coarser scale and only those fine tiles supported by local data, high prediction error, uncertainty, or an intervention query are instantiated.

Even the consensus policy does not demand agreement. For two views a and b that can be compared on support c , define $r_{ab}^{(c)} = \mathcal{T}_{a \rightarrow c} X_a - \mathcal{T}_{b \rightarrow c} X_b$. A compatibility pseudo-

likelihood may be

$$\Psi_{ab}^{(c)} \propto \exp \left[-\frac{1}{2} r_{ab}^{(c)\top} W_c r_{ab}^{(c)} \right],$$

where W_c depends on calibration, uncertainty, and whether the two sources are expected to measure the same process. Unless W_c and the residual law have been validated as a generative error model, this factor produces a generalized or pseudo-posterior, not a calibrated Bayesian posterior. Persistent disagreement can remain as a source-specific residual or branch the posterior into alternative latent explanations.

This resolution lattice is the interface for heterogeneous evidence. A high-density electrical source estimate, lower-resolution hemodynamic field, tract endpoint distribution, behavioral target, and focal stimulation field attach at the spatial and temporal supports they actually resolve. Gradients may cross a calibrated map, but no modality is upsampled into fictitious precision. The same primitive applies to cortical fields, nested connectome graphs, vascular and acoustic solvers, hippocampal episodes, temporal language structure, and regional task objectives. Figure 5 shows three common policies using two levels only for visual clarity; a compiled model may use an arbitrary and source-dependent set.

4.3 Recurrent Prediction, Denoising, and Causal Restraint

Predictive coding provides one useful cortical hypothesis: descending pathways convey context-dependent predictions, ascending pathways convey residual or evidence, and local recurrence estimates precision and resolves ambiguity (Bastos et al., 2012). Iterative cortical inference can also be implemented with denoising or diffusion-like updates over a structured latent canvas:

$$Z^{k+1} = Z^k + \eta_k s_\theta \left(Z^k, Y, C, \mathcal{G}, k \right) + \sigma_k \xi_k.$$

Here, the anatomical schema constrains which blocks can influence which other blocks during iterative inference. This creates an explicit interaction graph inside the algorithm. It does *not* imply that the biological cortex performs the reverse process of a contemporary diffusion transformer. Reverse diffusion, recurrent inference, energy minimization, predictive coding, and neural relaxation can generate related computational signatures while making different mechanistic commitments. SC-WBD keeps them as interchangeable backends whose causal predictions can be compared.

The distinction between model class and scientific claim is essential here. Structured denoising, predictive coding, recurrent state estimation, and energy-based relaxation may share an interaction topology and still imply different latent variables, update laws, laminar directions, timing, and perturbational responses. Efficiency or next-step prediction alone cannot select among them. Each backend must be evaluated against observations that are sensitive to its proposed mechanism and against equally expressive alternatives that share the same anatomy.

4.4 Neurodynamic Analogues for Machine-Learning Operators

Neuroscience-machine-learning correspondences are most useful when they specify a testable inductive bias and least useful when a familiar algorithm is declared to be what the brain is doing. Table 4 records the intended level of commitment. Convolutional receptive fields can improve agreement with aspects of human visual recognition (van Dyck et al., 2021), and goal-driven deep networks have predicted response structure in high-level visual cortex (Yamins

et al., 2014); however, recurrence remains necessary for challenging or temporally incomplete recognition (Tang et al., 2018; Kar et al., 2019). Sparse brain-derived topology has also improved the efficiency of some spiking reinforcement-learning models (Wang et al., 2024), but this establishes an engineering prior rather than functional homology.

4.5 Multirate Co-simulation

Electrical, cognitive, vascular, endocrine, learning, and structural states cannot be advanced efficiently at one clock. A multirate scheduler gives every field an update policy, interpolation contract, and error budget. Fast modules advance continuously or event by event; slow modules update only when their inputs materially change. Synchronization is forced when a fast transition changes the support of a slower process, such as sustained activity altering metabolic demand or a salient prediction error opening a plasticity window. Temporal coarsening contributes to reported posterior uncertainty rather than remaining an invisible implementation choice.

A learned propagator is not assumed to form a semigroup. For every permitted pair of steps, the scheduler measures

$$\epsilon_{\text{sg}}(\Delta_1, \Delta_2; x) = \left\| \Phi^{\Delta_1 + \Delta_2}(x) - \Phi^{\Delta_2} \left(\Phi^{\Delta_1}(x) \right) \right\|_W.$$

Adaptive stepping or substitution of one coarse update for several fine updates is disabled when this residual exceeds the declared validity tolerance. Below tolerance, its empirical distribution enters the numerical and model-discrepancy budget; it is not described only as temporal coarsening uncertainty.

5. Subcortical, Cerebellar, and Hippocampal Systems

SC-WBD does not implement one undifferentiated limbic module. The label mixes anatomical, historical, and functional categories at a granularity too coarse for intervention modeling. Hypothalamic and brainstem systems estimate and regulate bodily variables, organize autonomic and defensive responses, and provide arousal and homeostatic constraints. Amygdalar systems learn biological relevance and configure sensory, mnemonic, autonomic, and action preparation; they are not a scalar fear or valence node. Basal-ganglia loops participate in action selection, vigor, working-memory gating, and reinforcement-dependent policy. Thalamic nuclei contribute to routing, gain, synchrony, state control, and recurrent cortical computation. Cerebellar loops learn fast forward predictions and calibrated residual corrections across motor and non-motor domains.

These systems warrant more engineered regional backends than a generic transformer block because their circuit motifs, inputs, outputs, and physiological variables are comparatively constrained. Even here, the model retains competing hypotheses. Dopamine is not equated with reward, serotonin with punishment, acetylcholine with attention, or norepinephrine with arousal. Neuromodulators are receptor-, target-, timescale-, and state-dependent control fields that alter gain, plasticity, exploration, persistence, and uncertainty.

5.1 Hippocampal High-Dimensional Codes

The hippocampal-entorhinal system requires rapid binding, separation of similar episodes, completion from partial cues, sequence organization, and replay into cortex. Place and grid codes suggest spatial scaffolds that can be recruited for more general relational organization (Moser

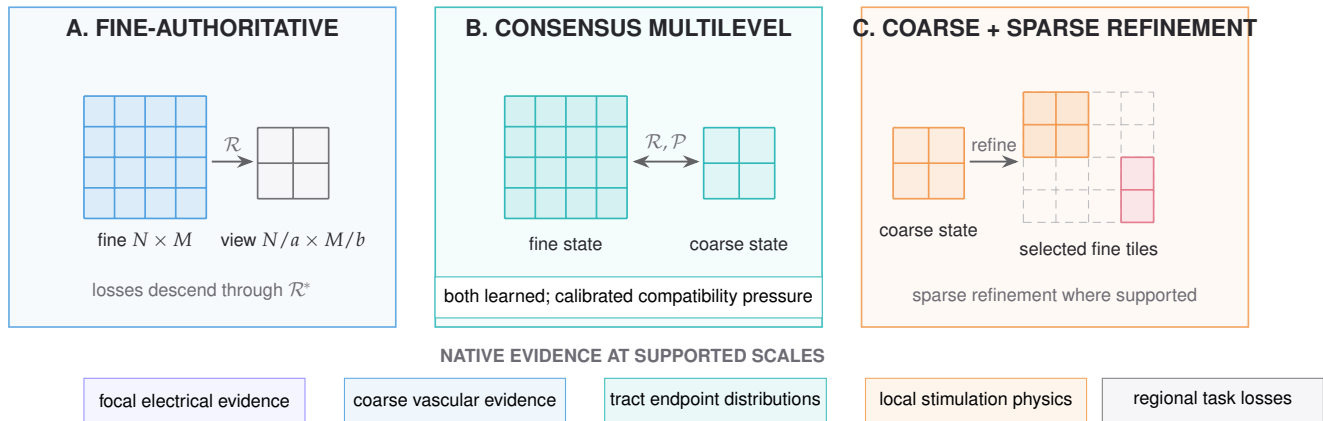


Figure 5 Three multiresolution state semantics. Resolution is not limited to preprocessing. A fine-authoritative field materializes derived coarse views; a consensus model gives multiple levels their own degrees of freedom under continuous agreement pressure; a coarse-authoritative model instantiates only selected high-resolution tiles. The schema declares the authority policy, projection operators, gradient routes, uncertainty, and refinement triggers. Every view and map carries separate bias and variance terms. Heterogeneous observations attach at the levels their physical point-spread functions justify. Two dyadic grids are drawn for clarity; actual views may use any number of non-dyadic, non-grid, and source-specific scales, and need not admit one consistent global field.

et al., 2015). The reference architecture therefore permits a high-dimensional sparse state

$$H_t = \{k_t, v_t, g_t, c_t, \rho_t\},$$

where k_t is a cue or index, v_t is bound content, g_t is a multiscale relational or grid-like code, c_t is temporal and contextual state, and ρ_t is retrieval confidence. Very tall vectors and similarity products are useful implementation devices, but a bare softmax retrieval equation does not distinguish hippocampal hypotheses. The claim is algorithmic, not cellular: neuronal populations need not literally execute a floating-point dot product. Attractor dynamics, dendritic nonlinearities, oscillatory phase, sparse expansion, and recurrent inhibition can implement selection and error correction through different mechanisms.

Alternative backends include modern Hopfield-style associative memory, Vector-HaSH-like factorized codes, spiking associative networks, successor representations, and complementary-learning-system models (Kumaran et al., 2016; Ravichandran et al., 2024). They share an interface but make different predictions about capacity as a function of code dimension and sparsity, interference as episodes accumulate, cue-degradation curves, theta/gamma timing, pattern separation, completion, replay order, and cortical consolidation. The benchmark therefore varies memory load, overlap, temporal lag, oscillatory phase, lesion/ablation, and replay opportunity. A hippocampal model is selected by those discriminating signatures rather than by whether its latent vectors are intuitively evocative.

6. Regional Pretraining and Whole-Brain Assembly

Training an undifferentiated whole-brain model end to end would permit high-capacity regions to absorb functions that should be localized elsewhere and would make failures difficult to interpret. Waiting for a homogeneous, simultaneous, whole-brain-whole-body dataset would create the opposite failure: most useful regional and behavioral evidence would be discarded because it lacks every modality. SC-WBD instead uses a staged, source-aware curriculum. Regional and motif models first learn from partial trajectories; typed interfaces make those lessons composable; connectome assembly then permits end-to-end correction without pretending that every datum supervised every state.

6.1 Stage I: Regional Phenotype Pretraining

Each regional family is pretrained on the measurements and tasks most informative about its function. Early visual fields receive retinotopic and natural-image dynamics, with gaze and pupil trajectories when available. Auditory fields receive spectrotemporal, scene, and speech tasks. Motor, somatosensory, spinal-interface, and cerebellar systems receive kinematic, proprioceptive, tactile, nociceptive, force, and error-correction trajectories. Hippocampal systems receive episodic, relational, navigation, replay, and interference objectives. Brainstem, hypothalamic, insular, and autonomic interfaces receive cardiac, respiratory, endocrine, digestive, temperature, pain, fatigue, and other interoceptive series. Cellular, laminar, receptor, spectral, and intrinsic-timescale priors regularize the state transition where they are empirically defensible.

Isolated pretraining does not imply that a region has one function or can be understood outside the network. Its purpose is to initialize a plausible operator family and prevent whole-brain optimization from beginning as a blank slate. A module is trained behind its declared ports using recorded, population-derived, or simulated boundary-state distributions. Boundary randomization, delayed and corrupted inputs, missing channels, and counterfactual interventions prevent it from depending on one artificial neighbor. Region-specific losses are retained after assembly so the whole model cannot silently repurpose a visual, cerebellar, or homeostatic module merely because another allocation improves a global metric.

6.2 Stage II: Interface and Pathway Calibration

Pairs and motifs are then trained through their declared interfaces: feedforward and feedback visual pathways, auditory-language loops, hippocampal-cortical replay, cortico-basal-ganglia-thalamo-cortical gating, cerebellar predictive loops, and interoceptive-affective control. Boundary adapters are penalized when they transmit information outside their declared type or exploit numerical side channels. Delay, direction, and topographic footprint are calibrated against physiology and perturbational data where available.

6.3 Stage III: Heterogeneous Sliced-Trajectory Training

A training episode need only illuminate a slice of the model. Let episode e declare an observed subgraph S_e , source-

Computational device	Neurodynamic analogue	ana-	Permitted interpretation in SC-WBD	Engineering status	Mechanistic evidence
Local convolution, mesh kernel, or neural field	Topographic receptive fields, short-range horizontal coupling, traveling waves	receptive	Strong prior for spatial locality; the learned kernel need not equal a synaptic circuit	High	Medium
Block-sparse attention or graph operator	Selective long-range corticocortical and corticosubcortical communication	long-range and com-	Useful topology and routing prior; attention coefficients are not tract strengths	High	Low
Recurrent predictive or error-correcting update	Reciprocal hierarchies, contextual prediction, mismatch responses	hierarchies, prediction,	Competing implementations must be distinguished by laminar, timing, and perturbation data	Medium	Medium
State-space model, reservoir, or neural ODE	Low-dimensional population trajectories and intrinsic regional timescales	pop-ulation regional	Compact dynamical surrogate; latent axes are not automatically biological variables	High	Low–medium
Modern Hopfield memory, key–value retrieval, or sparse high-dimensional code	Hippocampal pattern separation, attractor completion, relational and episodic indexing	pattern attractor relational	Similarity and error-correcting memory are plausible algorithmic correspondences; a dot product is not a literal neuronal operation	Medium	Low–medium
Mixture-of-experts routing and conditional computation	Thalamic routing, basal-ganglia gating, context-dependent cortical recruitment	basal-ganglia cortical	Models selective access and resource allocation without assigning one router to one nucleus	Medium	Low
Forward model and residual controller	Cerebellar prediction and error-dependent calibration across motor and cognitive loops	prediction error-dependent	Strong functional analogy in constrained domains; implementation remains multilevel	Medium	Medium
Diffusion or denoising trajectory	Iterative inference, relaxation, sampling, and pattern completion	inference, relaxation, sampling, and pattern completion	An algorithmic inference family, not a claim that cortex implements reverse diffusion	Medium	Low
Per-edge operator registry and differentiable composition	Distinct pathways and local circuits transform signals through heterogeneous biophysical mechanisms	transform signals through heterogeneous biophysical mechanisms	Scientific software interface for comparing ODE/PDE, field, SSM, spectral, attention, point-process, and surrogate operators; flexibility itself supplies no neural evidence	High	Neutral
Global latent workspace or broadcast token	Widespread availability, ignition, reportability, and flexible control	availability, reportability, and flexible control	Optional coarse-grained measurement plane; competing consciousness theories remain unresolved	Low–medium	Contested

Table 4 Machine-learning and neurodynamic analogues. Engineering status rates usefulness as an implementation prior; mechanistic evidence rates support for biological interpretation. The two ordinal judgments are not combined into one incommensurable confidence score.

native observations Y_e , optional interventions U_e , context C_e , coordinate and calibration record C_e , and observation operators O_e . A variational sliced-trajectory objective has the form

$$\mathcal{L}_{\text{slice}} = \sum_e \left\{ -\mathbb{E}_{q_e} \log p(Y_e | X_{S_e}, \theta_e, O_e, C_e) + \text{KL}[q_e(Z_e) \| p(Z_e | U_e, C_e)] + \mathcal{R}_{\text{ports}} + \mathcal{R}_{\text{dynamics}} + \mathcal{R}_{\text{cal}} \right\}, \quad (4)$$

where $Z_e = (X_{\partial S_e}, X_{\text{mis}}, \theta_e)$, $X_{\partial S_e}$ denotes uncertain boundary state, and X_{mis} denotes other marginalized state needed by the episode. The first two terms are a negative evidence lower bound only when the stated probability models are normalized and evaluable. Boundary, scale, distillation, and compatibility penalties remain auxiliary pseudo-losses; they are reported separately and do not inherit posterior-calibration claims from the generative likelihood. Unobserved interior modules are frozen, marginalized, or sampled from a population posterior according to whether gradients have a defensible causal path to them. Missing data are not imputed as labels. If an eye-tracking dataset says nothing about hypothalamus, the hypothalamic state is not trained to an arbitrary default; if a simulated spinal signal is uncertain, that uncertainty enters through the boundary distribution.

The source library is deliberately broad:

- **public research data** provide anatomy, electrophysiology, neuroimaging, behavior, tasks, and perturbations under reusable governance;
- **private and partnered data** add devices, clinical contexts, longitudinal outcomes, rare tasks, and high-value paired modalities while retaining access and use restrictions;
- **prospectively acquired data** are designed around identifiability—repeated calibration, synchronized clocks, multimodal tasks, state probes, and bounded interventions—rather than accumulated only because they are convenient;
- **simulated data** exercise physics, embodiment, missingness, causal interventions, and rare regimes, but remain simulator-conditioned evidence and cannot establish biological validity.

Useful slices range from narrow to richly embodied. Eye movements alone supervise oculomotor control and temporal policies; eye movements paired with visual content additionally supervise retinotopic sampling, salience, and world-centered stabilization. Simulated humanoids can provide synchronized motor commands and visual, auditory, tactile, proprioceptive, vestibular, nociceptive, and interoceptive consequences. Those streams are routed through the ports corresponding to peripheral receptors, spinal and brainstem pathways, thalamic relays, cerebellar loops, somatotopic cortex, and motor effectors rather than attached to a generic embodiment vector. Digestive, cardiorespiratory, autonomic, hormonal, and immune proxies can train

homeostatic interfaces even when no cortical recording accompanies them. Naturalistic audiovisual data, language, navigation, social interaction, sleep, learning, and rehabilitation trajectories constrain other overlapping slices.

Quarantined exploratory teacher prior. Reported inner experience may eventually motivate a deliberately low-authority distillation experiment, but it is not part of the core training path. A time-indexed report $R_e = \{(r_k, [a_k, b_k], q_k)\}_k$ records content, an uncertainty interval, confidence, and elicitation conditions. Candidate text, audio, or video renderings could be passed through a population encoder such as TRIBE v2 (d’Ascoli et al., 2026) and attached only to corresponding perceptual, language, or hemodynamic observation interfaces.

The chain is explicitly lossy: retrospective report $\rightarrow$ rendered proxy $\rightarrow$ perception-trained teacher $\rightarrow$ average-subject cortical prediction $\rightarrow$ subject registration. Imagery and perception differ in amplitude, extent, and hierarchical profile; population organization is not individual organization. The factor therefore supplies no new neural observation, receives no authority over latent mechanisms, subcortical or bodily state, and cannot pull an individualized posterior toward an average brain. It is a discrepancy-inflated interface regularizer, never a subject likelihood.

The experiment remains off by default. It is enabled only in a preregistered branch comparing no teacher, output and intermediate-feature distillation, matched generic multimodal features, generic smoothness, time-shuffled and report-mismatched teachers, and perception-versus-imagery calibration. It is retained only if it improves held-out measured subject-specific neural or behavioral forecasts beyond every matched control. Teacher agreement, fluent predictions, or smoother cortical maps are not validation. The long-term direction remains for SC-WBD itself to accept an image, sound, linguistic stream, or closed-loop context and emit individualized electrophysiological, hemodynamic, physiological, and behavioral predictions through native heads.

Heterogeneity creates strong confounding risks. Each episode retains participant, site, device, montage, preprocessing, task, session, body state, coordinate transforms, and selection policy. Source- and site-specific nuisance heads absorb stable instrument effects; hierarchical effects separate population, individual, and session variation; propensity or design factors are included when an intervention was not randomized; negative-transfer gates can stop an incompatible dataset from updating a shared operator. Gradient weights depend on estimated bias, variance, model discrepancy, and overlap in the target validity domain—not sample count alone. Cross-source agreement is evidence only after shared preprocessing and common-cause confounds have been examined.

Typed ports are the alignment substrate. A visual module learned from eye-plus-image data and an oculomotor module learned from kinematics can exchange only the documented retinal, gaze, efference-copy, and uncertainty variables. A motor controller learned in simulation may initialize a population prior, but measured human electrophysiology and behavior determine which parts survive transfer. This makes the coarse general model progressively useful before individual fine-tuning while preserving the ability to identify which source taught which parameter.

6.4 Stage IV: Connectome Assembly and Joint Multimodal Training

The compiler assembles motifs under the probabilistic connectome and trains the joint model on the full mixture of

sliced and synchronized sensory, motor, neural, physiological, behavioral, and interventional sequences. A fully paired multimodal recording is thus a particularly informative episode, not a requirement imposed on the corpus. A representative assembly objective is

$$\begin{aligned} \mathcal{L} = & \lambda_{\text{slice}} \mathcal{L}_{\text{slice}} + \lambda_{\text{obs}} \mathcal{L}_{\text{paired}} \\ & + \lambda_{\text{dyn}} \mathcal{L}_{\text{forecast}} \\ & + \lambda_{\text{int}} \mathcal{L}_{\text{perturb}} + \lambda_{\text{anat}} \mathcal{R}_{\text{topology}} \\ & + \lambda_{\text{scale}} \mathcal{R}_{\text{cross-scale}} + \lambda_{\text{homeo}} \mathcal{R}_{\text{stability}} \\ & + \lambda_{\text{cal}} \mathcal{L}_{\text{calibration}}. \end{aligned}$$

Losses are attached to declared roles and resolution levels rather than uniformly to every latent cell. Some states are directly reconstructed, others predict the future, others constrain conservation or homeostasis, and some are latent but tested through intervention. A coarse-modality residual may backpropagate through restriction and prolongation maps into a finer cortical field without claiming that the modality measured that field directly. Inverse-variance weighting is used only after modeled bias and shared confounding are considered; otherwise a precise but systematically distorted source could dominate the joint objective. End-to-end gradients may update microdynamics and boundary maps, while hard anatomical constraints remain compiled and soft constraints update only within their priors. Alternating curricula revisit regional and motif objectives between whole-system updates to detect functional reassignment, interface collapse, catastrophic forgetting, and simulator shortcuts.

6.5 Stage V: Individualization Without Catastrophic Identity Drift

Population parameters are adapted to the subject using anatomy, repeated recordings, tasks, and behavior. Stable subject parameters are separated from session state:

$$\theta_{p,s} = \mu + \alpha_{g(p)} + \delta_p + \zeta_{p,s},$$

with centered or sum-to-zero population effects $\sum_g n_g \alpha_g = 0$, hierarchical shrinkage $\delta_p \sim \mathcal{N}(0, \Sigma_{\text{person}})$, and $\zeta_{p,s} \sim \mathcal{N}(0, \Sigma_{\text{session}})$ or a declared temporal state model. Non-centered parameterizations may improve inference, but they do not remove the need for identifying constraints and recovery tests.

Sleep, fatigue, stress, medication, electrode impedance, and head pose should not be consolidated as permanent traits. Conversely, a stable preference or learned skill should not be erased when one session is atypical. Continual learning therefore uses replay, posterior regularization, explicit provenance, and uncertainty-triggered deferral before modifying slow subject parameters.

7. Multimodal Assimilation and Individualized Intervention

7.1 Inference Over State, Parameters, Topology, and Model Class

Given observations $Y_{1:M}$, interventions U , and context C , the system estimates

$$p(X_{0:T}, \theta_p, \mathcal{G}_p, \mathcal{M} \mid Y_{1:M}, U, C, A_p),$$

where $\mathcal{M}$ indexes alternative dynamical backends. This generative inversion is directly continuous with DCM’s program of inferring hidden neural state and effective coupling through modality-specific forward models (Friston et al., 2003); SC-WBD extends the software and data contract

rather than claiming that decomposition as new. Variational inference, sequential Monte Carlo, ensemble filtering, amortized posterior networks, neural-ODE initialization, and simulation-based inference can be combined according to the tractability of a subsystem (Kashyap et al., 2023). The posterior includes observation-model uncertainty: an EEG mismatch may arise from neural state, conductivity, electrode placement, or artifact, and the model should not force all error into cortex. Repeated-session and multi-device data are used to distinguish stable subject effects from source bias and session variance whenever the design permits that decomposition.

The conditioning set remains source indexed. Assimilation proposes factors between native views and latent fields; it does not overwrite the source data with a fused estimate. A query to read or write state specifies anatomical domain, scale, temporal window, clock, reference frame, calibration chain, units, support kernel, and tolerated uncertainty. EEG samples need not be downsampled to the fMRI repetition time, and fMRI voxels need not be assigned sensor-space electrical precision. A continuous-time or multirate posterior links them through their observation operators, timing uncertainty, and hemodynamic state. The runtime can answer from a native state, traverse a calibrated map, instantiate a local refinement, return a distribution over possible fine states, or report that the requested cross-scale relation is unresolved.

Cross-modal prediction reduces, but does not remove, non-identifiability. A trajectory constrained by EEG should predict held-out hemodynamic or behavioral consequences; anatomy constrained by diffusion MRI should predict plausible propagation delays; a state inferred from passive data should predict perturbational response. When two models remain observationally equivalent but imply different treatments, the output is an unresolved causal ambiguity, not an averaged recommendation.

7.2 TMS and Transcranial Focused Ultrasound

The TMS stack separates coil pose, induced field, cortical orientation, tissue coupling, immediate population response, network propagation, and plasticity. A target is not only a scalp coordinate but a tuple of subject surface, network, state, phase, task, waveform, dose, and learning window. Position in millimeters, roll, pitch, and yaw are recorded in the coil frame and transformed through tracker, fiducial-defined head, MRI, cortical-surface, and optional atlas frames. Orientation relative to the local cortical normal and tangents is retained. Neuronavigation data show that distance and angular placement errors materially alter the modeled cortical field (Caulfield et al., 2022); population atlas coordinates can also miss an individual functional target (Uehara et al., 2024). Connectivity-guided accelerated TMS results motivate individualized targeting (Cole et al., 2020), but they do not establish that a whole-brain model can prospectively select the best protocol for a new person. Pose accuracy, field accuracy, target engagement, network change, symptom change, and comparative clinical utility are separate validation levels.

The field solver writes at the mesh and temporal scale supported by device and tissue physics. Its effect may then be restricted into a coarser whole-brain state or trigger refinement of selected cortical tiles. The system does not paint a focal exposure directly onto a parcel scalar, nor does it materialize the entire cortex at field-solver resolution when only a bounded support is causally addressed.

The transcranial focused-ultrasound stack similarly separates tracked transducer pose, skull acoustics, steering commands, in situ pressure and thermal estimates, uncer-

tain cellular coupling, and downstream network dynamics. The planned focus and realized exposure are distinct random variables; tracking, registration, array steering, skull modeling, and coupling each contribute bias and variance (Phipps et al., 2024). Because low-intensity neuromodulatory mechanisms remain incompletely resolved, multiple tissue operators should coexist under posterior model comparison. Safety and exposure limits remain outside the learned therapeutic objective and follow independent biophysical standards (Aubry et al., 2025).

For a bidirectional non-invasive interface, read and write channels need not be the same method. tFUS may be evaluated as a write channel while EEG, MEG, fMRI, behavior, or another validated sensor supplies the read channel. Functional ultrasound imaging and transcranial ultrasonic stimulation should not be described as a mature unitary read-write technology merely because they share acoustic physics.

7.3 EEG Decoding, Control, and Cognitive Inputs

EEG can support constrained communication, motor imagery, attention classification, error detection, and computer control. Language priors can increase decoder fluency, but they can also conceal how little information came from the neural signal. SC-WBD therefore reports neural evidence separately from language-model completion and evaluates whether decoded content generalizes across prompts, sessions, vocabularies, and task structures. Non-invasive word decoding remains much less informative than unrestricted thought reading (d’Ascoli et al., 2025).

Images, sounds, language, social interaction, and tasks enter through perceptual and action models. Their effects depend on autobiographical memory, speaker and agent models, current affect, body state, expectation, and context. This is the practical basis of individualized cognitive intervention: generate or select an input expected to move a person’s latent state toward a bounded, measurable target, then update from their actual response. It is also why a generic prompt-to-outcome table is insufficient.

7.4 Risk-Sensitive Closed-Loop Control

Candidate protocols are selected across the posterior:

$$u^* \in \arg \max_{u \in \mathcal{A}_{\text{safe}}} \mathbb{E}_{p(X, \theta, \mathcal{M})} [B(X_T) - \lambda C(u) - \beta L_{\text{harm}}(X_T, u)] - \gamma \mathcal{U}_{\text{epi}}(u).$$

where B is predicted benefit, C is burden, and L_{harm} is an outcome loss averaged once under the posterior. $\mathcal{U}_{\text{epi}}$ separately penalizes reducible model uncertainty and disagreement; it does not double-count aleatoric outcome risk. Independently validated device, exposure, and protocol limits define the feasible set $\mathcal{A}_{\text{safe}}$ outside the learned objective. The controller may choose a safer measurement or reversible probe instead of an intervention. Figure 6 emphasizes that successful signal reconstruction is only the beginning of this loop.

8. From Closed-Loop Dynamics to World and Self Models

8.1 Environmental and Bodily Constraints on Latent Structure

The transition from neurodynamics to concepts need not be treated as a jump. An organism repeatedly encounters structured regularities: objects persist through occlusion, actions have delayed consequences, body states change the

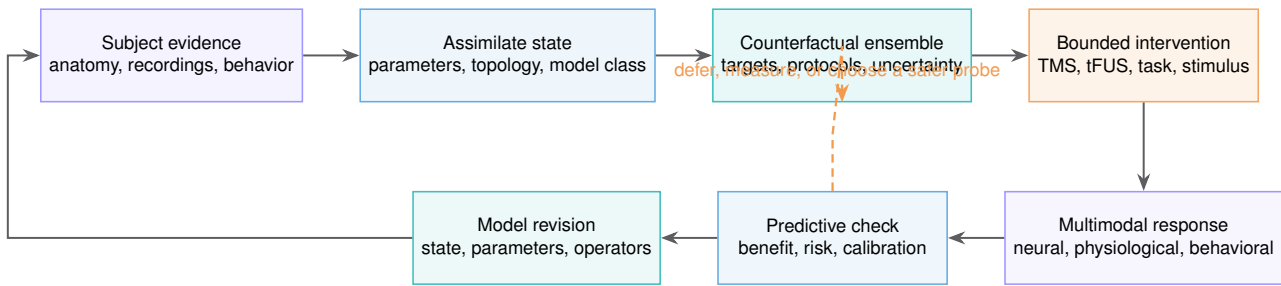


Figure 6 Closed-loop individualization. Measurements constrain a posterior over latent state, parameters, topology, and model class. Candidate actions are evaluated across posterior models, applied only through bounded protocols, measured multimodally, and used to revise the model. Clinical or wellness claims require prospective decision-level validation beyond successful signal reconstruction.

value of outcomes, and other agents respond contingently. A bounded nervous system cannot retain the full sensory history. To predict and control, it must compress trajectories into latent variables that preserve consequences relevant to future action. Craik's early formulation of an internal, small-scale model of external reality and the organism's possible actions already states this computational pressure (Craik, 1943); cognitive-map experiments established that behavior can reflect relational organization not reducible to a chain of rewarded responses (Tolman, 1948).

The important constraint is closed-loop alignment. A factor is useful when it continues to predict sensory change across viewpoints, actions, and time. Stable environmental causes therefore encourage increasingly world-centered latent structure; stable action-to-sensation and interoceptive contingencies encourage increasingly body- and self-centered structure. Active interaction is developmentally informative because the learner observes both a motor command and its lawful sensory consequence. Classic reafference theory distinguished self-generated from externally generated sensory input (von Holst and Mittelstaedt, 1950), and active-versus-passive rearing experiments showed that matched visual exposure is not equivalent to movement-contingent experience for visually guided development (Held and Hein, 1963). Later sensorimotor accounts emphasized perception as mastery of action-dependent sensory regularities rather than passive reconstruction (O'Regan and Noë, 2001).

SC-WBD does not claim that evolution or development optimized one modern machine-learning loss, nor does it simulate morphogenesis. It uses these findings as curriculum constraints. Latent factors should be rewarded for remaining predictive under action, viewpoint, delay, perturbation, and modality change. Models trained only to interpolate passive recordings can learn shortcuts that have no stable relation to the agent's world or body. The embodied and intervention-rich trajectory slices in Section 6 provide the tests that passive signal reconstruction cannot.

8.2 Reafference, Interoception, Agency, and the Minimal Self

A minimal self-model can be operationalized without presupposing phenomenal selfhood. The system estimates which body it controls, which sensory changes are consequences of its commands, which variables must be regulated for viability, where its perceptual perspective is located, and which memories belong to the current continuing controller. Efference copy and forward prediction help distinguish refferent from exafferent change; proprioceptive and vestibular loops stabilize body-centered coordinates; interoceptive prediction couples perception to homeostatic control; memory and action selection bind these estimates over longer time.

Predictive-processing and active-inference accounts provide one formal family for relating perception, action, and

bodily regulation (Friston, 2010). Interoceptive inference makes explicit that affect and embodied self-estimation depend on inferred causes of internal bodily signals (Seth, 2013); minimal-self models similarly connect ownership, agency, and multisensory integration to hierarchical prediction (Limanowski and Blankenburg, 2013). These are candidate computational organizations, not proof that one free-energy objective explains the self. SC-WBD permits comparator, control-theoretic, state-space, predictive-coding, and learned alternatives to share observation ports and face the same perturbations.

World and self variables are coupled but should not collapse. A reaching error may reflect an external target shift, a body-model error, fatigue, or an incorrect forward model. A social prediction may reflect the other person's state, the subject's projection, or a shared environment. The model therefore maintains explicit self, other, body, and world attribution probabilities, with provenance and calibration, rather than assigning every prediction error to one undifferentiated latent state.

8.3 Grounded Conceptual Primitives and Limited Symbolism

Repeated latent factors become candidate conceptual primitives when they are stable enough to support recognition, action, counterfactual prediction, memory retrieval, and communication across contexts. Spatial containment, continuity, controllability, agency, object persistence, approach, threat, effort, relief, and social reciprocity need not begin as discrete tokens. They can first exist as reproducible geometries, attractor families, transition operators, or control-relevant partitions of sensorimotor and interoceptive state. Categories and names later provide economical handles on those grounded regularities. This avoids the regress in which symbols are defined only by other symbols (Harnad, 1990); perceptual-symbol accounts likewise motivate conceptual content that reuses modality and action systems (Barsalou, 1999).

At the spatial and temporal granularity envisioned here, a universal ontology of beliefs, goals, and events may remove more information than it adds. Cortical and subcortical states are graded, distributed, metastable, context-sensitive, and only partly reportable. SC-WBD therefore does not place a general symbolic event graph above the brain. The declarative brain schema describes memory layout, anatomical priors, and communication constraints; it is not evidence that the subject represents those declarations symbolically. Discrete variables are used where the experiment supplies them: a presented word, chosen response, episode boundary, explicit rule, report, or clinician annotation. Elsewhere, structured continuous and event dynamics remain primary.

8.4 Candidate Planes for Conscious Access and Self-Modeling

Closed-loop prediction, grounded concepts, or a body model do not by themselves solve consciousness. They make several experimentally accessible planes less discontinuous with the rest of the architecture:

- **global availability** asks which contents become usable for flexible control, working memory, report, and cross-system coordination;
- **self-other attribution** estimates whether a state, action, or consequence belongs to the subject, another agent, or the environment;
- **attention and access schemas** summarize limited processing and make that summary available to prediction and control;
- **counterfactual agency** represents which variables are controllable and how imagined actions alter predicted outcomes;
- **metacognitive and report access** models which states influence confidence, error awareness, deliberate explanation, and language.

Global-workspace accounts motivate tests of broad availability and report (Baars, 1988; Dehaene and Naccache, 2001); attention-schema theory motivates a simplified model of attention used for control and attribution (Graziano and Webb, 2015). SC-WBD implements such proposals as competing learned coarse-grainings or probes over common lower-level dynamics. It can compare ignition, broadcast, recurrent access, self-model, and other summaries without making any one the system's privileged substance. Contemporary theories remain contested (Seth and Bayne, 2022); a candidate plane is retained only if it predicts discriminative timing, report, confidence, lesion, or stimulation effects beyond simpler latent summaries.

One preregistered contrast makes this requirement concrete. A workspace-availability plane and an equal-capacity local-recurrence control are fit without using the critical condition. Under matched TMS dose and baseline state, the workspace model must predict a later, distributed, report-linked propagation component with specified latency, topography, cross-area availability, and trial-wise relation to report or confidence; the control predicts local or task-limited recurrence without that broadcast. The test is failed if the signatures are chosen after seeing the propagation map, if both models absorb the effect through a generic residual, or if the added plane does not improve held-out prediction. This contrast tests one operational consequence of availability, not consciousness itself.

Integrated information theory is outside the initial implementation scope (Tononi et al., 2016). Its system-level Φ is not treated as estimable from the proposed non-invasive observations, and an approximate correlate would not adjudicate the theory. SC-WBD may retain perturbational complexity or integration summaries as descriptive observables, but it will not label them Φ or use them as a consciousness ground truth.

This establishes a hypothesis ladder rather than a derivation: closed-loop invariants may support body/world segregation; segregation may support agency and self/other attribution; distributed access may support flexible report and metacognition. None of those steps is asserted to be sufficient for phenomenology. The selection of maintained boundaries, recursive self-modeling, affective geometry, and pattern re-entry as candidate abstractions was partly prompted by a non-peer-reviewed conceptual manuscript (Valdez, 2026b). This is provenance rather than authority: those abstractions enter SC-WBD only through operational predictions and receive no evidential weight from their

source.

8.5 Affect and Computational Empathy

Affect is not one valence scalar or an isolated limbic activation. It can be modeled as a history-dependent geometry over interoceptive regulation, action readiness, uncertainty, memory, social appraisal, controllability, and prospective cost. Valence and arousal are useful projections, but their apparent values can depend on which dimensions and timescales are collapsed. The affective readout must therefore retain the source states and uncertainty from which it was derived.

Computational empathy is calibrated inference over another organism's latent dynamics,

$$p(X_{\text{other}}, \pi_{\text{other}} \mid O_{\text{other}}, C, \mathcal{M}_{\text{self}}),$$

under an explicit self-other boundary. It does not imply direct access to another person's experience. It estimates a distribution over likely needs, pain, beliefs, agency, appraisal, and action tendencies; forecasts behavior and report; updates from error; and records projection bias. This is the computational scaffolding by which social mammals use their own embodied models to understand others, expressed as a falsifiable forecasting and control problem rather than a sentiment label.

9. Language, Report, and Person-Specific Reconstruction

9.1 Language as a Grounded Read-Write Channel

Language enters after, and recursively reshapes, the latent structures just described. Words can index perceptual categories, actions, relations, interoceptive states, social roles, memories, and counterfactuals because those structures already constrain what the subject can discriminate, predict, and do. Linguistic instruction then permits new compositions and long-range credit assignment without replaying every grounding experience. Speech and text are therefore neither an astronomical addition to neurodynamics nor the substrate of all thought: they are high-bandwidth, socially standardized read-write channels over partially grounded world, body, self, and other models.

A generic language model estimates a distribution such as $p(W_n \mid W_{<n}, C)$. It can imitate a person's style when supplied with biographical context, yet its hidden causal organization remains mostly that of the generic model. A person-specific brain model instead defines a joint process:

$$p_{\theta_p}(X_{0:T}, Y_{0:T}^{\text{neural}}, B_{0:T}, W_{1:N} \mid S_{0:T}, U_{0:T}, A_p, H_p),$$

where language W is emitted by, and feeds back into, the same latent state that generates neural signals, bodily response, action, memory retrieval, affect, and learning. Word choice, latency, hesitation, revision, omission, confidence, prosody, and the downstream effect of one's own utterance become observations of a continuing person model. Next-word models already explain variance in human language responses (Goldstein et al., 2022; Schrimpf et al., 2021); the stronger SC-WBD claim must be tested by cross-predicting nonlinguistic state and preserving individual causal regularities across contexts.

Language remains an unusually dense behavioral measurement, not a complete readout of consciousness, implicit skill, interoception, or memory. During training, generic language priors are separated from subject evidence. A fluent completion that is weakly constrained by measured neural

or behavioral state must be labeled as model completion rather than decoded thought or personal prediction.

9.2 Attractor, Pattern, and Predictive Fidelity

A stable person model can be evaluated at several non-equivalent levels. An *attractor model* reproduces characteristic basins, transitions, and recovery after perturbation. A *relational pattern model* preserves invariants among memories, values, affective regimes, concepts, and action policies across tasks or implementations. A *predictive person model* prospectively reproduces choices, reports, learning, and responses to unfamiliar interventions. The distinctions are operational definitions in this manuscript, not deductions from a named philosophical framework.

SC-WBD operationalizes the three fidelities separately:

- **attractor fidelity** is tested by perturb-and-recovery trajectories and characteristic switching among neural, affective, and behavioral regimes;
- **pattern fidelity** is tested by cross-modal, cross-task, and cross-backend invariants with explicit tolerance and uncertainty;
- **predictive fidelity** is tested prospectively on unfamiliar situations, learning episodes, social exchanges, and bounded interventions.

None establishes that experience is identical to a cause-effect structure, that a copied pattern is the same person, or that re-instantiation continues a first-person point of view. Those are philosophical hypotheses outside what neural resemblance, self-report, or predictive accuracy alone can settle.

9.3 A Reconstruction Ladder

Claims about a digital twin are divided into four levels:

1. **bounded behavioral reconstruction** predicts speech, choices, or actions in a stated domain and time horizon;
2. **multimodal structural reconstruction** preserves internal relations that jointly explain neural, physiological, and behavioral data;
3. **causal-dynamic reconstruction** reproduces recovery, learning, counterfactual judgment, and response to held-out intervention;
4. **phenomenological continuation** claims continuity of the experiencer rather than predictive similarity of a model.

The first three admit scientific and engineering tests. The fourth is not established by success at the first three. SC-WBD therefore uses *digital reconstruction*, *person model*, or *delegated behavioral model* unless a stronger claim becomes scientifically justified.

A delegated model estimates

$$p(a_p, w_p \mid s, h_p, v_p, X_p),$$

where s is the situation, h_p is relevant history, v_p represents inferred values and constraints, and X_p is estimated state. Every output distinguishes direct subject evidence, population priors, generic completion, and out-of-distribution extrapolation. A repeated individual pattern is not interchangeable with a plausible response from a generic agent. Authorization, domain limits, uncertainty, deferral, auditability, correction, and revocation are therefore primary model capabilities rather than product annotations.

10. Execution Architecture and Research Artifacts

The development stack is separated into interoperable packages:

- **schema and compiler**: anatomical entities, state types, geometry, resolution posets, authority policies, clocks, topology, operator dispatch, provenance, and uncertainty;
- **dynamics**: cortical fields, laminar and spiking models, hippocampal memory, subcortical controllers, hemodynamics, plasticity, and learned surrogates;
- **observe**: EEG, MEG, MRI-derived, fMRI, fNIRS, physiological, behavioral, and report forward models;
- **intervene**: electromagnetic, acoustic, sensory, cognitive, task, neurofeedback, and bounded control operators;
- **infer**: filtering, smoothing, simulation-based inference, model comparison, calibration, counterfactual ensembles, and control;
- **bench**: synthetic fixtures, unit and solver tests, reference datasets, ablations, posterior predictive checks, bias-variance ledgers, subgroup calibration, and audit reports.

Every module declares input and output spaces, units, coordinate frames, temporal support, differentiability, solver requirements, numerical tolerance, uncertainty, and evidence. Scientific schemas remain separable from particular implementations: laboratories can compare alternative thalamocortical, hippocampal, vascular, or language modules on identical anatomical priors, tasks, observations, and interventions.

Research artifacts are intended for public release without making openness the paper's scientific claim. The planned repository is [JacobFV/integrated-whole-brain-modeling-across-modalities-scales-and-dynamics](https://github.com/JacobFV/integrated-whole-brain-modeling-across-modalities-scales-and-dynamics). The initial public surface should include schemas, compiler code, synthetic data, reference configurations, benchmark definitions, and reproducible ablations. Each work package also receives a machine-readable claim manifest: thesis and schema versions, target claim, permitted source cards, baselines, acceptance thresholds, refusal fixtures, and explicit non-goals.

The development study targets 23 adult participants across 89 longitudinal sessions as a precision-neuroscience deep-sampling design (Poldrack et al., 2015; Gordon et al., 2017). The individual and their future sessions are the unit of inference; these totals are not a power claim for clinical efficacy and will be revised by prospective reliability and information-gain analyses. They are planning numbers, not completed recording. Any release of participant data depends on prospective ethics approval, consent that explicitly covers public distribution, de-identification and re-identification-risk review, and continuing governance. Neural and anatomical data cannot be assumed anonymous merely because names are removed.

11. Validation and Falsification

The first gate is the linear-Gaussian native-clock benchmark in Section 0.3, not a full neural simulator. Its Fisher-information and recovery results decide which coupling, delay, observation, and calibration parameters the next model is allowed to claim. Every later empirical subsystem repeats local structural identifiability, practical identifiability, prior-sensitivity, and synthetic recovery analyses; a parameter that is prior-dominated is reported as such rather than stabilized by an unacknowledged regularizer.

11.1 Numerical, Representational, and Physical Tests

The compiler is tested for shape, unit, coordinate, delay, and mask correctness. Modules must pass solver-convergence, stability, conservation, random-seed, and boundary-consistency tests where applicable. Electromagnetic and acoustic solvers are validated independently of neural-response models. Fine and coarse regional back-

ends must agree within a declared tolerance on boundary observables before adaptive resolution is used for inference.

11.2 Signal and Cross-Modal Tests

Held-out temporal, spectral, spatial, cross-frequency, state-transition, and behavioral statistics are tested beyond the fitting loss. Latent state inferred from one modality must predict another, and performance must be compared with simple autoregressive, dense neural, population, and subject-specific statistical baselines. A personalized model must outperform a population model on a new session; including the person's scan is not itself individualization. Metrics are reported with sampling variance, bootstrap or posterior intervals, estimated optimism from model selection, and bias analyses across session, device, acquisition site, anatomy, demographic strata, and task context. Aggregate accuracy cannot substitute for calibration within the intended deployment population.

Comparability of counts is not comparability of populations. Normalizing a log loss per observation makes designs that consume different numbers of observations arithmetically comparable; it does not make them scientifically comparable, because the mean is then taken over different observation sets. A design that additionally consumes a small set of poorly fit rows will show a large apparent degradation that is pooling arithmetic and not negative transfer, and the pooled figure will move in whichever direction the extra rows happen to push it. Designs consuming different observation sets must therefore be compared on their common set, with the design-specific rows reported separately and never averaged into the shared comparison. This is a measured near-miss in our own first-gate artifact rather than a hypothetical (Section 11.5), and it carries a second lesson that became visible only once the first was repaired: restricting the comparison to the common population removed the pooling defect and revealed that none of the fits being compared had converged. A held-out comparison is evidence about the designs only when it is taken over a common population *and* between fits that reached their optima. Either condition alone is insufficient, and a comparison that fails the second one is reporting the optimizer.

11.3 Perturbational and Decision Tests

The strongest causal tests hold out interventions and ask for direction, latency, spatial propagation, dose response, state dependence, and behavioral consequence. Decision validity is stricter: model-guided selection must improve prospective outcomes relative to existing targeting or control procedures. Until then, TMS, tFUS, wellness, and treatment outputs are research hypotheses, not clinical recommendations.

11.4 Required Ablations

At minimum, each claim should survive comparisons among:

- structured regional state versus one scalar or pooled vector per region;
- hybrid local field plus sparse long-range graph versus fully dense, graph-only, and uniformly convolutional models;
- single-resolution versus simultaneous fine/coarse cortical pyramids, arbitrary source-native resolution lattices, and sparse adaptive refinement, including scale- and parameter-matched controls;
- hard, soft, learned, randomized, and distance-matched topology;
- anatomically typed operators versus an equal-parameter generic operator;
- region-specific pretraining versus end-to-end blank-slate

training;

- population, session-adapted, and longitudinal subject models;
- language-only behavioral imitation versus a language process coupled to neural, bodily, memory, and action predictions;
- when the quarantined report/teacher experiment is enabled: no teacher, matched generic features and smoothness, shuffled/mismatched report, and perception-versus-imagery domain-shift controls;
- correlation fitting versus held-out perturbational prediction.

Topology is supported only if it improves data efficiency, calibration, out-of-distribution behavior, or causal prediction—not merely because it reduces parameters. A mechanistic module is supported only if removing or replacing it worsens a prediction uniquely associated with its mechanism. Every ablation reports both variance and plausible systematic error; a lower variance model is not preferred when it achieves stability by smoothing away the effect of interest.

11.5 First-Gate Results: the Section 0.3 Artifacts

Both Section 0.3 artifacts have been built and run. What follows are results about a *design*, obtained in a synthetic linear-Gaussian system whose truth is known by construction. They are methodological findings, not empirical neuroscience: no claim about any brain, any participant, or any recorded dataset is asserted here or anywhere else in this document. The distinction is the point of the exercise, and it is preserved deliberately.

Verdict. The benchmark reports INCOMPLETE under the decision rule fixed in its manifest before the run. Three criteria were fully evaluated and failed: fusion information, native-beats-resampled, and energy-matched intervention information. Two recovery criteria—nominal interval coverage and recovery improvement—were withheld as *not evaluated* rather than recorded as failures, under a convergence gate described below. The gate is disclosed as post hoc: the manifest fixed the five criteria before the run and did not contain it.

The fusion claim narrows to a claim about clocks—and the summary statistic conceals a decomposition that matters. Across all three regimes the joint native-clock design and EEG alone agree on the preregistered information statistic to six significant figures; the ratios are 1.0000004–1.0000858 (Table 5). Read carelessly that says “the hemodynamic modality is uninformative,” which is not what was measured. Two separate things are true, and the headline number is the product of both.

First, the reported quantity is a *minimum* eigenvalue, so it reports the worst-identified direction, and in every regime that direction is the conduction delay. Per parameter, the hemodynamic channel carries 0.1% to 1.0% of the electrical channel's Fisher information about the three coupling gains—small, real, and non-zero—and between 4×10^{-7} and 1×10^{-6} of it about the delay. In the reference regime the diagonal terms are 212.2, 57.2 and 16.0 against 0.85, 0.35 and 0.14 for the couplings, and 65.4 against 6.0×10^{-5} for τ . The delay result is structural rather than a matter of degree: a 12 ms conduction delay does not survive a multi-second hemodynamic convolution sampled once per second. The defensible statement is therefore that **the hemodynamic modality cannot improve the parameter that is hardest to identify**—not that it is uninformative, and a reader must not be permitted to make that substitution.

Second, the small coupling contribution that does exist is consumed by the hemodynamic observation model. Before

Held-out regime	best single modality	joint / best single	impulse, unmatched	impulse, energy-matched	naive resampling, (T4) / estimator
reference: 12 ms delay, prior-mean coupling	16.01 (EEG)	1.0000015	9.31×	0.839×	0.1218 / 0
weak coupling, 17 ms delay	1.840 (EEG)	1.0000004	27.91×	0.839×	0.01569 / 0
low SNR, 8.5 ms delay	13.74 (EEG)	1.0000858	6.94×	1.059×	0.03941 / 0

Table 5 Section 0.3 benchmark, three held-out regimes. The statistic is the minimum eigenvalue of the likelihood-only expected Fisher information over the preregistered subset $(a_{21}, a_{32}, a_{13}, \tau)$ after profiling out the observation nuisances, in the prior-standardized basis; the prior contribution is excluded. Columns 3–5 are ratios to the joint native-clock design. The last column gives the resampled design's value under (T4) and under the coarse model the resampling estimator actually uses: information present in the data, and information the estimator can use. Twenty-four replicates per regime.

profiling, the joint design's θ -block minimum eigenvalue exceeds EEG alone's by 0.3% to 0.9% (16.170 against 16.027 in the reference regime). Profiling out only the electrical nuisances leaves that gain intact (16.151 against 16.008); profiling out the three hemodynamic nuisances—the response shape, the undershoot, and the BOLD gain—removes it entirely, returning exactly EEG alone's value. The second modality spends its information estimating its own observation model from the same record. That diagnosis identifies a fixable half, and this program already names the remedy: the calibration-source register in the implementation supplement lists breath-hold/hypercapnia cerebrovascular-reactivity, arterial-spin-labeling and calibrated-fMRI acquisitions as supplying subject and session hemodynamic parameters, with uncertainty, for neural-to-BOLD observation models. Constraining those externally would stop the second modality paying for them out of the same data. The ceiling on that repair is the pre-profiling gain: about 0.9% on this statistic. Real, worth acquiring, not transformative—and it does nothing whatever for the delay.

Neither result is an artifact of the block-diagonal form of (T4): the exact joint information, obtained by whitening the stacked record, differs from the (T4) value by at most 0.2% here, so the near-identity survives the correction.

Naive resampling, by contrast, is structurally broken, and it is broken in a way that has to be stated precisely. On the one-second lattice $\partial\mu/\partial\tau \equiv 0$; the coarse model has no delay line, so the delay is not merely poorly estimated but unrepresentable. Under the estimator's own model the likelihood rank is 3 of 9, the profile rank over the preregistered subset is 1 of 4, and the profile minimum eigenvalue is exactly zero in every regime. Under (T4) the *same design* has likelihood rank 9 of 9 and profile rank 4 of 4. Both statements are true and they are about different objects: the information is present in the data, and the resampling estimator discards it. Collapsing the two views inverts the finding. The consequence is visible in recovery: in the reference regime the resampled design's empirical interval coverage for the three coupling gains is 0.33, 0.46 and 0.71 against a nominal 0.95.

The supported claim is therefore that **source-native clock handling beats resampling**, and that **adding a modality is not automatically beneficial**. These are different claims and only the first has support here. Table 1, row 1, is narrowed accordingly.

A calibrated impulse's apparent gain is its input energy. Compared without an energy control the impulse looks decisive: $6.9\times$ to $27.9\times$ the joint native-clock profile information, for an increase in total input energy of about 11%. With the background drive scaled so that total input energy matches the impulse-free design exactly, the ratios fall to 1.06, 0.84 and 0.84—below one in two of the three regimes. In those two regimes, spending the same energy on background drive buys *more* information about the preregistered subset than spending it on the impulse. An unmatched impulse comparison measures the energy budget, not the

perturbation, which is why the energy-matched design is required rather than supplementary and why the criterion was preregistered against the matched design.

Two design faults, self-identified and reported rather than absorbed. First, the reference regime places the true conduction delay at 12 ms, which is the prior mean. A design that learns *nothing* about the delay leaves it at the prior mean and so scores a perfect delay error: the resampled design, whose delay information is exactly zero, records a delay RMSE of 0 ms in that regime. Delay evidence from the reference regime is not discriminating, is flagged as degenerate in the artifact rather than averaged in, and the two held-out regimes place the delay away from the prior mean for precisely this reason. Second, the recovery estimator—a fixed-budget damped Newton run from the prior mean—did not converge in two of the three regimes, with median Newton decrements of 20.9 and 9.2–22.7 posterior standard deviations against a gate of 2.0. Coverage and recovery are properties of that estimator, so the two recovery criteria are withheld rather than reported as negative results; an unconverged sweep must not be allowed to masquerade as a falsification.

The second artifact. The end-to-end synthetic slice compiles the same system from a source schema under known transform bias, shared session calibration error, missing windows, unequal supports, and a deliberately misspecified residual. It refused four of four invalid schemas while accepting the valid one; produced a leakage-free grouped split and fired on an ungrouped positive control, which is what distinguishes a leakage audit from a leakage assertion; passed subgroup calibration; and detected the misspecified module, an AR(1) colored residual on one hemodynamic parcel. Detection was sensitive rather than specific: seven of fourteen module tests were flagged at $\alpha = 0.05$, of which six were false positives. One criterion failed and one could not be evaluated. Recovery interval coverage was 0.16, 0.19, 0.53 and 0.50 against a nominal 0.95 over 32 simulated subjects; coverage is a property of the estimator as run, so that is a measurement. The held-out log-loss criterion is not. Its first form is the near-miss described in Section 11.2: the joint and single-modality designs were compared over different observation populations—845,280 rows against 844,800, the difference being 480 hemodynamic rows—so the pooled per-observation means were never comparable to begin with. Recomputing on the common 844,800 rows removes that defect and exposes a second, independent one. A comparison between two fits measures the *designs* only if both fits reached their optima, and none of the four did: the median Newton decrements are 12,238 (joint native-clock), 409 (EEG alone), 100,478 (fMRI alone) and 53,675 (naive resampling) posterior standard deviations, against the same gate of 2.0 used above, and the fMRI-alone fit never left the prior mean at all—zero drift in every parameter. The criterion is therefore reported as *not evaluated*: not passed, not failed, and no direction of effect asserted in either sense.

No per-observation loss from these fits is quoted here, because each of them is a property of an optimizer that had not arrived.

That is two criteria now withheld for non-convergence, across both artifacts and in the same run. Together they state something sharper than either alone: **the current limit on what this benchmark can conclude is the estimator, not the design.** Until the optimizer reaches its optima the recovery and held-out prediction questions are simply unavailable, and no amount of further modeling reaches them.

What this licenses. It licenses the narrowing above, and nothing about brains. The synthetic system is linear, Gaussian, three-region, and its truth is known; dose and state dependence are unavailable in it by construction, and its delay result is simulation recovery rather than held-out perturbation. It also asks the hemodynamic modality an electrophysiological question, for the reason given in Section 0.3: a system with no spatial structure and a parameter set of coupling gains plus one delay has nowhere for that modality's actual strengths to appear. "Cross-method integration is not worthwhile" does not follow from these numbers and is not asserted. Section 0.3 stated in advance that if the calculation did not go the claim's way, "the compiler may still be useful as a provenance system, but the claim that cross-method integration resolves those dynamics must be narrowed." That condition obtained, and that is the sentence now in force.

11.6 Decorative Guards, and Thresholds Without a Reference Class

The most transferable result of building the above was not a number. It was a class of defect that a validation program of this kind produces continuously and by construction, and which no amount of care alone removes.

A guard is decorative when its reading is constant with respect to the question it is asked. It still emits output, the output still looks like evidence, and it cannot come out any other way—which makes it worse than no instrument, because confident reasoning follows from it. A register kept during this build accumulated more than twenty such instruments in a single development period. Representative cases: gradient-permission patterns that were matched against parameter names a graph compiler rewrites, so the permission check was vacuous on the accelerator while every host-side test passed; a memory cap that accounted for host pages while the allocation was on the device, reporting 8.17 GB against 97.9 GB actually held; a leakage audit that existed, raised correctly when called, and was never called by the trainer, while a schema field asserted that it had passed; a stop threshold that read identically whether its hypothesis was true or false; and a normalizer-selection metric on which one candidate scored exactly 1.00 at every percentile because the statistic equals one by algebra for that candidate.

The general test costs nothing and should be applied to every guard, gate, and acceptance criterion in this program:

For each guard, name a state of the world in which it reads differently. If there is none, it is decoration.

Two corollaries were expensive enough to be worth stating separately. First, an *absent* record is indistinguishable from a clean one: a field written only on success is not a record, so the null case must write something explicit rather than nothing. Several of the cases above have that shape. Second, a verification is evidence only about the path it exercises; a check that reaches its subject by a different route than production does can pass while production fails. The remedy is to make the guard fail on purpose—break the thing

deliberately and confirm the alarm sounds—rather than to observe that it is green.

That four of the recorded instances were found *inside* machinery built to catch exactly this failure is not an embarrassment to report reluctantly. A falsification apparatus that detects its own defects at that rate is the best available evidence that it is measuring rather than performing, and the register is retained for that reason.

A summary statistic can strip the qualifier that made it true. This is the same pattern applied to reporting rather than to instruments, and it occurred while preparing Section 11.5. The statistic there is a *minimum* eigenvalue; carried one step away from its definition it becomes "the second modality adds nothing," which is a different and false proposition, because what was measured is that the second modality cannot improve the worst-identified direction while it does inform the others. Nothing in the number is wrong and no instrument misread; the compression is what fails, and it fails quietly because the compressed form is shorter, more quotable, and still numerically supported by the figure it came from. That is the warning of Section 13 about redescription, operating on this project's own results rather than on a theory. The test is not whether a summary is defensible but whether it establishes the proposition it will be read as establishing—so where a single number aggregates directions that behave differently, the decomposition must travel with it.

A preregistered threshold with no reference class is a guess with a timestamp. This is the sharpest instance, and it is a limit of preregistration as a technique rather than a lapse in applying it. A training acceptance condition—running-minimum simulated forecast negative log likelihood below 1.0 by step 900—was fixed in the repository before the data existed, never moved, honored to the letter when it resolved at 1.1200, escalated the moment a sibling clause fired, and adjudicated by a party who neither set it nor trained the model. Every procedural safeguard the project has was applied, correctly, in order. The result is still uninterpretable, because no capacity-matched baseline and no prior measurement existed that could say whether < 1.0 was ever achievable for a 1.76M-parameter model on 37 simulated shards with roughly 40% of its regional-timescale prior either never delivered to the backend or clamped to the support boundary. Without a reference class the threshold cannot distinguish a model that underperformed from a number that was never achievable: the two produce an identical reading.

Two consequences follow for the acceptance criteria in Sections 11.1–11.4. First, **process discipline cannot manufacture a reference class**; preregistration fixes *when* a commitment is made and does nothing about *what* is committed to, so a threshold must additionally be shown capable of reading differently in the world where its hypothesis is false. Second, a preregistration inherits the defects of the instrument it was calibrated against, and freezing it makes that inheritance harder to see rather than easier; when the instrument is later found defective the commitment becomes *uninterpretable* rather than wrong, and the honest response is to report it as uninterpretable, in separate layers—the literal fact, the interpretation, and an adjudication made by someone who is not its author—rather than to re-set it once the trajectory is visible.

The remedy in force is therefore **matched controls rather than absolute thresholds**. A control encodes no assumption about what is achievable, and both sides of the comparison move together under an instrument rescale, so the comparison survives a correction that would render an absolute bar meaningless. Where this program states an acceptance criterion as an absolute number, that number must be accompa-

nied by the reference class it was set from, or replaced by a control.

11.7 Second-Gate Results: a Trained Whole-Brain Attempt, and What It Measured

A conditional multirate model over 414 parcels was compiled, trained on open electrophysiology, and scored on a participant-disjoint held-out split against five statistical baselines. It lost to all of them. The result is reported here because the diagnosis, not the loss, is the finding: **the artifact was the equal-capacity generic-operator control of Section 11.4's first required ablation, and the comparison it lost was not matched on four of the eight stages separating latent state from reported scalar.**

The artifact was the control, not the model. Section 11.4 opens with *structured regional state versus one scalar or pooled vector per region*. The trained system was the second of these: one operator string applied to every region, a dense state tensor of identical width at every parcel, and the hippocampal, subcortical and multiresolution machinery implemented but never instantiated. Its designation asserted the first. The interface specification that produced this narrowed $X_i \in \mathcal{X}_i$ — a state *space* indexed by region — to a uniform per-parcel feature vector, and did not record that a narrowing had occurred. The controlling defect is therefore the *undeclared* narrowing: a stated one is a decision a review process can attack, while an unstated one is invisible to a process built out of attacking stated things. The remedy in force is a register of declared narrowings in which any divergence from this document is enumerated with its reason, and anything absent from the register is a defect by definition.

The loss was in the variance channel, and the conditional mean was never the problem. Decomposed on a ladder that holds each arm's conditional mean fixed and varies only the variance model — identically across arms, and validated by reproducing every baseline's own calibration to four decimals — the penalty resolved as: scale 0.4467 nats, channel 0.1113, horizon 0.0096, state 0.19–0.26. The cause was a single scalar. The observation head's log-variance parameter was flat to 3% across channels and asserted a predictive variance of 1.31 against a realised residual variance of 3.97: uniformly overconfident by $3.0\times$. Its optimum has a closed form; it was trainable in one late curriculum stage that ran 900 steps in 134 seconds, and reached 20% of the way. A second head's noise parameter never received a gradient at all. With per-window errors retained and paired participant-clustered intervals computed, the same artifact *beats every baseline on the conditional mean*, persistence by $-3.20 [-3.94, -2.51]$ and the nearest competitor by $-0.103 [-0.214, -0.003]$.

Two conclusions hold simultaneously and neither rescues the other. The comparison was not calibration-matched: five baselines received held-out per-horizon, per-channel variance calibration and the model received none, and the two arms lacking held-out calibration are exactly the two with positive excess over the Gaussian-entropy floor. And the model is contracted to carry $X_i^{\text{uncertainty}}$ as regional state and emitted a constant, which is a defect no instrument change repairs.

Capacity matching covers two of eight stages. A comparison between two arms passes through inputs, conditioning, state, observation interface, head parameterisation, score, split, and optimiser. Parameter and compute

budgets constrain the third and the eighth. Four defects found in this program sit on the intervening four and none is a budget: an observation interface that silently narrowed one arm's mean path to two exported dimensions against the other's eighteen; a variance parameter with no path from state; a score computed under calibration given to one side only; and a per-subject baseline reduced to its global fallback by a participant-disjoint split, while reporting 71 fitted models in use when the true count was zero. Each was filed separately and they are one class. A matched budget with an unmatched interface is an unmatched comparison wearing a passing check.

A verification apparatus inherits its surrogate's assumptions as blind spots. The identifiability machinery of Section 11.5 detected most defects in this program and could not have detected this one. Its exact Fisher computations run on a linear-Gaussian surrogate in which state-independent innovation covariance is a *theorem* — it is why a single Riccati recursion can be shared across the trajectory — so a constant log-variance is correct there and the failure is not merely undetected but unrepresentable. Its remaining arms measure parameter intervals, not predictive intervals, and the failure was in the predictive channel. Every check was green and every check was correct. The question to ask of a guard is therefore not only whether it can fire, but whether the failure it targets is *representable in the model it runs against*.

The ablation could not have seen its own effect. Minimum detectable effect, derived from the six paired participant-clustered comparisons the run itself produced, is 0.1404 nats at 27 held-out participants. The matched-calibration gain is 0.1025–0.1249 and sits below it; the horizon term at 0.0096 is fifteen-fold below it. The corpus held 109 participants, so the test fraction was a choice rather than a constraint, and at the full split the limit falls to 0.0699. An inconclusive result under the original fold would have been indistinguishable from a true null. Power is a property of the design, not a diagnosis available after the fact, and it is now fixed before any arm trains.

What the successor run may claim. Measured separation among cortical parcels supports a binary partition and nothing finer: under a spin-based null with FDR control, a seven-network candidate separates 6 of 21 pairs and a unimodal/association split separates its single pair. Cytoarchitectonic classes fail globally on every measured block available and are therefore carried as description and barred from justifying a partition. Because that separation is measured on receptor profile, intrinsic timescale, and myelin — properties an operator's *parameters* carry — the licensed claim is narrowed in advance to *a two-family cortical state partition beats a uniform one at matched capacity*, and the operator-typing ablation is held out as a separate comparison the present evidence cannot resolve. Fourteen subcortical parcels across seven families are declared out of claim: no family-level effect there is measurable at any participant count this corpus supports.

11.8 Third-Gate Results: What a Second Artifact Reveals

A second whole-brain model was compiled and trained on the anatomical prior described in Section 3, with heterogeneous per-family regional state stored in a padded layout with enforced per-family spans. (An earlier draft of this sentence said *segment rather than padded*; that is the alternative layout, which is built and tested but is not what these

weights use. The artifact records layout: `family_padded`, and 47.34% of the 414×59 state plane is pad, because two hippocampal parcels of width two set the width for every region.) Final scores are not reported here; the run is incomplete at the time of writing and any number would be an interim read. What is reportable, and is the point of this subsection, is a class of defect that only a *second* artifact can expose.

What the artifact actually is, and how we found out.

The run trained on *simulated* trajectories alone. Its configuration renamed all six curriculum stages; six gates inside the trainer select behaviour by matching the *previous* run's stage names, and five of the six therefore returned the wrong answer—no gradient was ever taken on measured recordings, the per-stage gradient allowlist fell back to a wildcard, boundary randomisation and haemodynamic state were disabled in the rollout, and no individualiser was constructed, so the stage named for individualisation ran ordinary simulator training. The sixth gate, the only reason the run trained at all, is correct by accident: it is the single gate written as an inequality, and an unrecognised name happens to satisfy it. Nothing raised an error, the objective decreased throughout, and every monitored quantity stayed nominal for nine hours. We therefore describe the checkpoint as a **simulation-trained model over a measured anatomical prior, evaluated on recordings it never saw**—a transfer result rather than a held-out one.

The mechanism generalises past this codebase. Each gate is a lookup keyed on a name, with a default that grants rather than refuses: `PERMISSIONS.get(name, ("*",))` and `name in (...)` are *unfalsifiable by construction*, in that no name exists which they reject. An unwired stage and a fully-permitted stage produce identical readings.

Two features of the episode bear on how such defects are found at all. First, it surfaced while reading a code path *adjacent* to the one under investigation—a branch in the evaluation code that had never executed—rather than from any monitor. Second, and more uncomfortably, a complete and tested remedy already existed in the repository, written eleven hours before the run began, together with six failing tests naming the defect precisely; the patch was never applied and the failures were never read. A red test that nobody reads is worse than no test, because it converts “*there is no check for this*”, which prompts action, into “*the suite is somewhat red*”, which does not.

Six defects, one geometry. Seven distinct faults stood between a configured run and its first prediction. Six of them shared a single shape: **a code path declared for two model arms and exercised on only one**. The comparison arm carries no engineered regional backends and no family layout, so it silently skips every branch that exists to serve the other; the automated tests also ran on CPU, while the artifact ships on an accelerator. Both axes drifted, and the faults lived in the gap—in the trainer, in a regulariser, in the serving path, in the intervention path, and in the evaluator.

Two of the six would not have raised an exception. Without binding the conditioning parameters, the engineered subcortical backends run on their defaults: the anatomical conditioning that motivates per-family backends is discarded and *the run completes*. And a regulariser that reads an attribute absent on the heterogeneous arm would, had it resolved to anything benign, have been missing from precisely the arm whose mechanistic backends it exists to protect—leaving the learned residual free to absorb the dynamics, so that an ablation would measure a residual block wearing a mechanistic model's name. These do not

produce a wrong number. They produce a right-looking number from a model that has quietly lost the mechanism the experiment exists to test.

Naming is a second such class. A refusal that forbids emitting a checkpoint under a designation it has not earned was already in force, mutation-tested, and firing. It guarded nothing here, because every naming fault found was in a path it does not watch: an evaluation stamped the *previous* model's identifier onto its own results and ranked its own arm under that label, in the very file a public model card reads its scores from; and checkpoint discovery takes a directory name as the designation without consulting the identifier recorded inside. The refusal guards *emission*; these were *writing* and *discovery*. A guard on one verb is not a guard on a noun.

Worse, the second naming fault survived the fix for the first, because the fix was verified by searching for the designation as spelled in the first—and the second was spelled differently. A literal repaired by search is repaired only in the spelling one thought of. The durable remedy is to derive the name once and let every consumer read it, leaving no spelling to miss, with a fallback that is deliberately *unnamed* rather than any real designation: an unnamed artifact is a visible defect, a misnamed one is not.

The uncomfortable inference. All of these faults were as old as the code. None was detectable while the project had a single artifact, because every name is trivially correct when there is only one thing to name, and no arm can diverge from another when there is only one arm. A programme that intends to compare models should build its second model earlier than feels necessary—not for the comparison, but because the second artifact is the instrument that makes the first one's defects visible.

An instrument that cannot reach a failure is not evidence about it. A density collapse in the amortized posterior was diagnosed and repaired through three sequential causes. Four isolation probes were built to reproduce it—randomised data, real corpus data, progressive masking, and a masked-batch probe constructed specifically to mimic the sliced-trajectory regime. All four were stable at exactly the configuration that diverged in the training loop, and the last compared the candidate repair against the status quo and found them indistinguishable. The repair was nevertheless correct: the production run held the posterior flat where every prior attempt had diverged by two to three orders of magnitude at the same step. The probes were not weak evidence against the repair; they were no evidence at all. This is the converse of the surrogate blindness recorded in Section 11.5—there a failure was *unrepresentable* in the test model, here it was representable but *unreachable*—and both are cases of a green instrument saying nothing while appearing to say something reassuring.

12. Confidence-Calibrated Development Horizons

12.1 Near-Term: Executable Multimodal Subsystems

The most credible initial deliverable is not an accurate simulation of every mental event. It is a compiler and model family that can represent structured regional state, reproduce selected EEG/fMRI/behavioral dynamics, perform cross-modal forecasts, and expose topology and uncertainty. Constrained EEG control, error-aware decoding, and offline comparison of TMS target hypotheses are compatible with

current methods. Performance should be expected to be domain specific and session sensitive.

12.2 Intermediate: Prospective Individualized Control

A more demanding milestone is prospective prediction of a person's response to a bounded perturbation across sessions. This requires reliable state estimation, device and tissue models, longitudinal calibration, and decision comparators. Success in one indication or protocol would not validate a general brain twin. The appropriate claim would be that a particular model improves target selection or control for a specified population and endpoint.

12.3 Long-Term: Person Models and Generated Cognitive Interventions

Generating language, images, sounds, or interaction policies for an individualized cognitive objective requires a stable model of interpretation, memory, affect, agency, and social context. A system could eventually forecast how a person is likely to understand an explanation or respond to a sensory program, and could update from their measured response. A delegated behavioral model could become useful in bounded situations. Neither capability would by itself demonstrate unrestricted mind reading, a complete psychological model, consciousness, or survival of personal identity.

13. Limitations and Epistemic Boundaries

The architecture can fail by becoming a comprehensive vocabulary with little identifiable content. Heterogeneous modules may exchange variables with the same name but incompatible semantics. Learned residuals may absorb errors that should falsify anatomy or physiology. Multi-modal heads may share confounds. Increasing biological detail may improve credibility while worsening parameter identification. Compute-efficient surrogates may preserve ordinary trajectories while failing exactly under the perturbations for which the model is needed.

The mature connectome is neither fully known nor fully fixed. Human tractography lacks reliable direction for many pathways and contains false positives and false negatives. Functional connectivity is dynamic and method-dependent. Local cortical circuitry, receptor states, and neuromodulatory effects cannot be inferred uniquely from ordinary non-invasive recordings. The hard/soft/proposed edge scheme manages these limitations; it does not eliminate them.

Predictive processing, diffusion, global workspaces, affective manifolds, and high-dimensional hippocampal codes can each become redescrptions that explain every outcome after the fact. Their value depends on preregistered contrasts that make different timing, laminar, generalization, lesion, or stimulation predictions. The same discipline applies to digital reconstruction: fluent self-description can be generated without causal fidelity, and causal fidelity would still not settle phenomenal continuity.

Finally, individualized neurotechnology is high stakes. Safety limits, clinical oversight, informed consent, privacy, equitable evaluation, and regulatory requirements cannot be learned away by a model. Wellness framing does not weaken the evidential obligation when an intervention can alter neural function.

14. Discussion

SC-WBD's central proposal is a change in the unit of whole-brain modeling. Instead of a matrix of scalar nodes, it uses an operator graph over heterogeneous structured state spaces. Instead of selecting either local convolutional organization or long-range connectome structure, it compiles both. Instead of treating one resolution as the brain, it treats resolution as a query-dependent and potentially simultaneous scale space with tested boundary maps. Instead of calling every learned dependency causal, it separates structural priors, effective operators, functional statistics, observations, and interventions. Each of those objects retains its own estimated bias, variance decomposition, provenance, and validity domain.

Anatomy can compile into a sparse interaction grammar that makes every latent block and permitted pathway inspectable while leaving within-region computation and message content learnable. The grammar is nevertheless probabilistic and revisable: biology requires uncertain edges, local fields, recurrent loops, state-dependent routing, multirate adaptation, and slow structural change. Operator dispatch supplies a controlled model comparison, not an exemption from measuring which dynamics are actually useful.

One of these commitments has already been made smaller by measurement rather than by argument. The Section 0.3 laboratory was built to ask whether native-clock fusion or a calibrated intervention increases likelihood information for a preregistered parameter subset. What it vindicates is the treatment of clocks, not the addition of modalities: resampling onto a common lattice destroys delay information the data still contain, while the second modality adds about one percent on the coupling gains—and that percent is consumed estimating its own observation model—and nothing measurable on the delay, and a calibrated impulse's apparent advantage is accounted for by its input energy (Section 11.5). Section 0.3 wrote the consequence down in advance—"the compiler may still be useful as a provenance system, but the claim that cross-method integration resolves those dynamics must be narrowed"—and that consequence is now in force. The narrowing is the falsification program working, not a repair applied to it after the fact; and it is a result about a synthetic design, which leaves every empirical question in this paper exactly where it was.

The architecture also narrows the role of symbolism. Declarative schema is essential for engineering and interpretability; broad symbolic cognition is not assumed to be the brain's operative format. Coarse consciousness-oriented planes are retained only where they connect distributed state to report, global availability, agency, affect, and testable behavior. Language is more central, not because thought is language, but because a longitudinal language stream is a dense causal trace of memory, appraisal, self-model, and social action when it is generated by the same person-specific dynamics.

This establishes a continuous research path from present neurotechnology to more ambitious person models. At every step, the relevant question is not whether the simulation looks brain-like, but whether its declared structure improves calibrated prediction under held-out conditions and interventions. If it does, anatomy and physiology have provided useful computational priors. If it does not, the compiler makes the failed commitments identifiable enough to revise.

15. Conclusion

SC-WBD is a technical thesis with a staged implementation and falsification program for integrated whole-brain model-

ing across modalities, scales, and dynamics. Its first claim-bearing deliverable is deliberately not a whole brain: it is a native-clock identifiability and synthetic-recovery laboratory that can show whether EEG-like dynamics, fMRI-like hemodynamics, and a bounded intervention jointly identify selected coupling and delay parameters. Failure narrows the fusion claim before architectural breadth can conceal it. That laboratory has now been run in the linear-Gaussian case, and it did narrow the claim: to source-native clock handling, with the modality-fusion and intervention criteria unsupported at the scope tested (Section 11.5). This remains a result about a design. No empirical claim about brains is asserted anywhere in this document.

The intended whole-brain model's regions contain structured fields, tensors, graphs, event processes, or associative states; its edges are typed operators rather than scalar weights; and its adult anatomical scaffold is strong but uncertain and slowly adaptive. A compiler translates this schema into local cortical neighborhoods, simultaneous multiresolution pyramids, sparse long-range communication, multirate schedules, observation models, intervention ports, losses, bias-variance ledgers, and auditable state layouts.

The immediate payoff is cross-method inference without false equivalence. Individual MRI geometry can constrain where and through what tissue an effect is expressed; EEG can constrain rapid latent trajectories that fMRI cannot resolve; fMRI can constrain distributed and slow consequences that sensor-space EEG cannot localize uniquely; and calibrated TMS or tFUS fields can test which candidate dynamics respond causally. Coordinate chains, spatial support, temporal support, and bias-variance ledgers make these methods mutually informative while preserving disagreement and source-native resolution.

The framework supports regional pretraining, heterogeneous sliced-trajectory learning, motif and whole-brain assembly, individualized posterior inference, adaptive resolution, and explicit comparison of mechanistic and learned backends. It does not require a homogeneous whole-brain dataset before useful training begins. It limits symbolic modeling to justified coarse-grainings, derives language interfaces only after grounded world and self dynamics, and distinguishes behavioral reconstruction from identity or phenomenal continuation. Its value will depend on ablation, cross-modal generalization, identifiability, held-out perturbation, prospective decision tests, and calibrated refusal when evidence is insufficient. The implementation contract makes every expansion name its sources, transforms, gradient permissions, uncertainty status, baselines, and disabling result. Those requirements turn an ambitious whole-brain position from an aesthetic synthesis into an executable scientific program.

Appendix. Candidate Evidence Program by Source Role

The development corpus is heterogeneous by design. The table below summarizes the source roles required to rough out SC-WBD before individual fine-tuning. The full dataset-and-calibration register is distributed as the companion implementation supplement; that register is a discovery and planning surface, not a claim that every resource is accessible, licensed, harmonized, or useful. Each admitted snapshot requires a source card, version and hashes, parent lineage, governance decision, native support, transform/clock manifest, bias-variance-discrepancy status, gradient permissions, and fixed split role.

The prospective 23-adult, 89-session design is intended as longitudinal precision neuroscience in the lineage of MyConnectome and the Midnight Scan Club, where repeated within-person measurement supports individual inference rather than population-level clinical efficacy (Poldrack et al., 2015; Gordon et al., 2017). Session allocation should be chosen by prospective information gain and reliability analysis, not defended by the nominal totals alone. Early sessions prioritize geometry, clock, lead-field, field, baseline-state, and repeatability calibration; later sessions hold out tasks and bounded perturbations for prospective validation.

An optional population teacher such as TRIBE v2 may be tested only as a weak distillation prior at visual, auditory, language, or hemodynamic observation interfaces (d'Ascoli et al., 2026). Retrospective report, rendered proxy, perception-trained teacher, average cortical response, and subject registration form a lossy domain-shift chain; imagery and perception must not be assumed to share response amplitude or hierarchy. The teacher receives no authority over subject-specific latent dynamics, subcortical state, bodily state, or mechanistic parameters and is removed unless it improves held-out empirical prediction beyond matched generic-feature and smoothing controls.

Data and Materials Availability

No participant dataset is analyzed or claimed in this thesis. The development plan targets a precision-neuroscience deep-sampling design of 23 adults across 89 longitudinal sessions; the individual, rather than a pooled group mean, is the principal unit of inference. Enrollment, sample counts, and data release remain prospective and require ethics approval, explicit sharing consent, de-identification, and re-identification-risk review. Schemas, synthetic fixtures, benchmark definitions, compiler code, and governance-permitted research artifacts are planned for github.com/JacobFV/integrated-whole-brain-modeling-across-modalities-scales-and-dynamics. The comprehensive candidate-source register is supplied separately as an implementation supplement; inclusion in that register is not a claim of access or fitness.

Author Contributions

Jacob Valdez: conceptualization, architecture, methodology, software design, visualization, project formulation, and writing.

Funding

No dedicated external funding is reported for this development manuscript.

Conflict of Interest

The author is Founding Architect at SuperCognition Labs, the organization proposing and developing SC-WBD, and may have future commercial interests in neurotechnology applications. No clinical capability, regulatory clearance, therapeutic efficacy, or completed human-data release is asserted.

Generative AI Statement

Generative AI assisted with drafting, language editing, and the production of schematic vector figures. Scientific claims, citations, and any future submitted version would require independent author verification under the applicable journal policy.

Document Status

This is a technical thesis and development contract: draft, pending development. It is not a clinical protocol, regulatory submission, completed software release, or report of empirical findings. Its scientific claims are conditional on the falsification gates stated herein.

References

- Albers, K. J., Liptrot, M. G., Ambrosen, K. S., Røge, R., Herlau, T., Andersen, K. W., et al. (2022). Uncovering cortical units of processing from multi-layered connectomes. *Frontiers in Neuroscience* 16, 836259. doi:10.3389/fnins.2022.836259
- Aqil, M., Atasoy, S., Kringelbach, M. L., and Hindriks, R. (2021). Graph neural fields: A framework for spatiotemporal dynamical models on the human connectome. *PLOS Computational Biology* 17, e1008310. doi:10.1371/journal.pcbi.1008310
- Aubry, J.-F., Attali, D., Schafer, M., Fouragnan, E., Caskey, C., Chen, R., et al. (2025). ITRUSST consensus on biophysical safety for transcranial ultrasound stimulation. *Brain Stimulation* 18, 1896–1905. doi:10.1016/j.brs.2025.10.007
- Baars, B. J. (1988). *A Cognitive Theory of Consciousness* (Cambridge: Cambridge University Press)
- Barsalou, L. W. (1999). Perceptual symbol systems. *Behavioral and Brain Sciences* 22, 577–660. doi:10.1017/S0140525X99002149
- Bastos, A. M., Usrey, W. M., Adams, R. A., Mangun, G. R., Fries, P., and Friston, K. J. (2012). Canonical microcircuits for predictive coding. *Neuron* 76, 695–711. doi:10.1016/j.neuron.2012.10.038
- Brynjarsdóttir, J. and O'Hagan, A. (2014). Learning about physical parameters: The importance of model discrepancy. *Inverse Problems* 30, 114007. doi:10.1088/0266-5611/30/11/114007
- Cakan, C., Jajcay, N., and Obermayer, K. (2023). neurolib: A simulation framework for whole-brain neural mass modeling. *Cognitive Computation* 15, 1132–1152. doi:10.1007/s12559-021-09931-9
- Cannon, R. C., Gleeson, P., Crook, S., Ganapathy, G., Marin, B., Piasini, E., et al. (2014). LEMS: A language for expressing complex biological models in concise and hierarchical form and its use in underpinning NeuroML 2. *Frontiers in Neuroinformatics* 8, 79. doi:10.3389/fninf.2014.00079

Evidence role	Representative source families	Permitted training or calibration use	Non-claim / dominant failure
Adult anatomy and phenotype	HCP, UK Biobank, BigBrain, Julich-Brain, Allen Human Brain Atlas, receptor and myelin atlases	Population priors over surfaces, parcels, laminae, tracts, tissue fields, regional cell/transcript/receptor phenotypes, and their covariance	Atlas frequency is not subject truth; post-mortem, registration, age, ancestry, acquisition, and tractography biases remain explicit
Fast neural dynamics	OpenNeuro/BIDS EEG and MEG, MOABB, Temple EEG, Cam-CAN, OMEGA, intracranial clinical archives, Allen/CRCNS electrophysiology	Native-clock spectral, transient, propagation, state-transition, and regional interface likelihoods; local micro/mesoscale priors where species and preparation permit	Sensor mixing, artifacts, clinical sampling, medication, species, and electrode-placement bias prevent direct whole-brain labels
Hemodynamic and metabolic dynamics	HCP fMRI, Midnight Scan Club, MyConnectome, Natural Scenes Dataset, StudyForrest, Narratives, BOLD5000, fNIRS and calibrated-physiology studies	Slow distributed observation heads, neurovascular parameter priors, naturalistic task trajectories, and repeated-person calibration	BOLD is not neural ground truth; timing, vascular, motion, task, and preprocessing discrepancies are modeled
Structural connectivity and material calibration	HCP diffusion, TractoInferno, ISMRM tractography challenge, histological tract atlases, IT'IS tissue properties, conductivity, dielectric, acoustic, thermal, and MRI-relaxometry references	Topology priors, tract endpoint/delay distributions, uncertainty benchmarks, and coefficients for EEG/MEG, TMS, tES, MRI, tFUS, and thermal forward models	Streamlines are not axons or effective weights; ex-vivo tables and solver agreement do not establish in-vivo target engagement
Perturbation and intervention	TMS-EEG/TMS-fMRI studies, motor and cognitive perturbations, tES, DBS/iEEG stimulation, tFUS EEG/evoked-response datasets, sensory and pharmacological challenges	Direction, latency, dose, state dependence, field-to-neural response, and causal model discrimination under explicit device and sham models	Requested dose is not delivered field; confounding, blinding, pose, artifact, carryover, and mechanism uncertainty can invalidate causal labels
Embodied boundary dynamics	Eye tracking with visual content, motion capture, speech/audio conversation, touch and pain, vestibular/proprioceptive tasks, ECG/PPG/EDA/respiration, electrogastronomy, endocrine/metabolic studies, wearables, simulated humanoids	Train sensorimotor, spinal/brainstem, oculomotor, autonomic, interoceptive, homeostatic, social, language, and environmental ports without requiring a simultaneous brain recording	Boundary supervision does not identify unobserved neural implementation; construct labels, simulator gap, synchronization, and demographic bias remain
Behavior, report, and longitudinal person evidence	Navigation, learning, sleep, decision, affect, language, interaction, daily-life sampling, repeated deep-phenotyping, and prospectively designed SC-WBD sessions	Policy/readout likelihoods, session-versus-trait hierarchy, stable preference and skill dynamics, uncertainty-aware report models, and future-session tests	Behavioral fit is not neural localization, consciousness, identity, or clinical utility; selection, self-report, and language priors may dominate
Simulation, teacher, and negative controls	Electromagnetic, acoustic, hemodynamic, and body simulators; embodied agents; null generators, perturbed mechanisms, and optional population encoders	Exercise interfaces, rare regimes, identifiability, missingness, and distillation at authorized observation boundaries; create falsification controls before empirical training	Synthetic rows are never independent people. Teacher agreement is not evidence; imagery/perception and population/individual domain shifts require quarantine and empirical ablation

Table 6 Evidence program summarized by epistemic role. Sources are admitted to update named modules, ports, scales, and nuisance parameters rather than a monolithic whole-brain target. A long source list is valuable only when each family has a defined gradient path, exclusion test, and held-out claim.

- Caulfield, K. A., Fleischmann, H. H., Cox, C. E., Wolf, J. P., George, M. S., and McTeague, L. M. (2022). Neuronavigation maximizes accuracy and precision in TMS positioning: Evidence from 11,230 distance, angle, and electric field modeling measurements. *Brain Stimulation* 15, 1192–1205. doi:10.1016/j.brs.2022.08.013
- Cole, E. J., Stimpson, K. H., Bentzley, B. S., et al. (2020). Stanford accelerated intelligent neuromodulation therapy for treatment-resistant depression. *American Journal of Psychiatry* 177, 716–726. doi:10.1176/appi.ajp.2019.19070720
- Craik, K. J. W. (1943). *The Nature of Explanation* (Cambridge: Cambridge University Press)
- d'Ascoli, S., Bel, C., Rapin, J., Banville, H., Benchetrit, Y., Pallier, C., et al. (2025). Towards decoding individual words from non-invasive brain recordings. *Nature Communications* 16, 10521. doi:10.1038/s41467-025-65499-0
- d'Ascoli, S., Rapin, J., Benchetrit, Y., Brooks, T., Begany, K., Raugel, J., et al. (2026). A foundation model of vision, audition, and language for in-silico neuroscience. *arXiv preprint arXiv:2605.04326* doi:10.48550/arXiv.2605.04326
- Davison, A. P., Br"uderle, D., Eppler, J., Kremkow, J., Muller, E., Pecevski, D., et al. (2009). PyNN: A common interface for neuronal network simulators. *Frontiers in Neuroinformatics* 2, 11. doi:10.3389/neuro.11.011.2008
- Dehaene, S. and Naccache, L. (2001). Towards a cognitive neuroscience of consciousness: Basic evidence and a workspace framework. *Cognition* 79, 1–37. doi:10.1016/S0010-0277(00)00123-2
- Friston, K. (2010). The free-energy principle: A unified brain theory? *Nature Reviews Neuroscience* 11, 127–138. doi:10.1038/nrn2787
- Friston, K. J., Harrison, L., and Penny, W. (2003). Dynamic causal modelling. *NeuroImage* 19, 1273–1302. doi:10.1016/S1053-8119(03)00202-7
- Galappaththige, S., Gray, R. A., Costa, C. M., Niederer, S., and Pathmanathan, P. (2022). Credibility assessment of patient-specific computational modeling using patient-specific cardiac modeling as an exemplar. *PLOS Computational Biology* 18, e1010541. doi:10.1371/journal.pcbi.1010541
- Glasser, M. F., Coalson, T. S., Robinson, E. C., et al. (2016). A multi-modal parcellation of human cerebral cortex. *Nature* 536, 171–178. doi:10.1038/nature18933
- Goldstein, A., Zada, Z., Buchnik, E., et al. (2022). Shared computational principles for language processing in humans and deep language models. *Nature Neuroscience* 25, 369–380. doi:10.1038/s41593-022-01026-4
- Gordon, E. M., Laumann, T. O., Gilmore, A. W., Newbold, D. J., Greene, D. J., Berg, J. J., et al. (2017). Precision functional mapping of individual human brains. *Neuron* 95, 791–807.e7. doi:10.1016/j.neuron.2017.07.011

- Graziano, M. S. A. and Webb, T. W. (2015). The attention schema theory: A mechanistic account of subjective awareness. *Frontiers in Psychology* 6, 500. doi:10.3389/fpsyg.2015.00500
- Harnad, S. (1990). The symbol grounding problem. *Physica D: Nonlinear Phenomena* 42, 335–346. doi:10.1016/0167-2789(90)90087-6
- Held, R. and Hein, A. (1963). Movement-produced stimulation in the development of visually guided behavior. *Journal of Comparative and Physiological Psychology* 56, 872–876. doi:10.1037/h0040546
- Henson, R. N., Abdulrahman, H., Flandin, G., and Litvak, V. (2019). Multimodal integration of M/EEG and f/MRI data in SPM12. *Frontiers in Neuroscience* 13, 300. doi:10.3389/fnins.2019.00300
- Jirsa, V. K., Proix, T., Perdakis, D., Woodman, M. M., Wang, H., Gonzalez-Martinez, J., et al. (2017). The virtual epileptic patient: Individualized whole-brain models of epilepsy spread. *NeuroImage* 145, 377–388. doi:10.1016/j.neuroimage.2016.04.049
- Kar, K., Kubilius, J., Schmidt, K., Issa, E. B., and DiCarlo, J. J. (2019). Evidence that recurrent circuits are critical to the ventral stream's execution of core object recognition behavior. *Nature Neuroscience* 22, 974–983. doi:10.1038/s41593-019-0392-5
- Kashyap, A., Plis, S., Ritter, P., and Keilholz, S. (2023). A deep learning approach to estimating initial conditions of brain network models in reference to measured fMRI data. *Frontiers in Neuroscience* 17, 1159914. doi:10.3389/fnins.2023.1159914
- Kennedy, M. C. and O'Hagan, A. (2001). Bayesian calibration of computer models. *Journal of the Royal Statistical Society: Series B (Statistical Methodology)* 63, 425–464. doi:10.1111/1467-9868.00294
- Kobeleva, X., López-González, A., Kringelbach, M. L., and Deco, G. (2021). Revealing the relevant spatiotemporal scale underlying whole-brain dynamics. *Frontiers in Neuroscience* 15, 715861. doi:10.3389/fnins.2021.715861
- Kumaran, D., Hassabis, D., and McClelland, J. L. (2016). What learning systems do intelligent agents need? complementary learning systems theory updated. *Trends in Cognitive Sciences* 20, 512–534. doi:10.1016/j.tics.2016.05.004
- Limanowski, J. and Blankenburg, F. (2013). Minimal self-models and the free energy principle. *Frontiers in Human Neuroscience* 7, 547. doi:10.3389/fnhum.2013.00547
- Luo, L. (2021). Architectures of neuronal circuits. *Science* 373, eabg7285. doi:10.1126/science.abg7285
- Markov, N. T., Ercsey-Ravasz, M. M., Ribeiro Gomes, A. R., et al. (2014). A weighted and directed interareal connectivity matrix for macaque cerebral cortex. *Cerebral Cortex* 24, 17–36. doi:10.1093/cercor/bhs270
- Moser, M.-B., Rowland, D. C., and Moser, E. I. (2015). Place cells, grid cells, and memory. *Cold Spring Harbor Perspectives in Biology* 7, a021808. doi:10.1101/cshperspect.a021808
- Oh, S. W., Harris, J. A., Ng, L., et al. (2014). A mesoscale connectome of the mouse brain. *Nature* 508, 207–214. doi:10.1038/nature13186
- O'Regan, J. K. and Noë, A. (2001). A sensorimotor account of vision and visual consciousness. *Behavioral and Brain Sciences* 24, 939–973. doi:10.1017/S0140525X01000115
- Phipps, M. A., Manuel, T. J., Sigona, M. K., et al. (2024). Practical targeting errors during optically tracked transcranial focused ultrasound using MR-ARFI and array-based steering. *IEEE Transactions on Biomedical Engineering* 71, 2740–2748. doi:10.1109/TBME.2024.3391383
- Poldrack, R. A., Laumann, T. O., Koyejo, O., Gregory, B., Hover, A., Chen, M.-Y., et al. (2015). Long-term neural and physiological phenotyping of a single human. *Nature Communications* 6, 8885. doi:10.1038/ncomms9885
- Ravichandran, N., Lansner, A., and Herman, P. (2024). Spiking representation learning for associative memories. *Frontiers in Neuroscience* 18, 1439414. doi:10.3389/fnins.2024.1439414
- Sanz-Leon, P., Knock, S. A., Spiegler, A., and Jirsa, V. K. (2015). Mathematical framework for large-scale brain network modeling in the virtual brain. *NeuroImage* 111, 385–430. doi:10.1016/j.neuroimage.2015.01.002
- Sanz Leon, P., Knock, S. A., Woodman, M. M., Domide, L., Mersmann, J., McIntosh, A. R., et al. (2013). The virtual brain: A simulator of primate brain network dynamics. *Frontiers in Neuroinformatics* 7, 10. doi:10.3389/fninf.2013.00010
- Schrimpf, M., Blank, I. A., Tuckute, G., Kauf, C., Hosseini, E. A., Kanwisher, N., et al. (2021). The neural architecture of language: Integrative modeling converges on predictive processing. *Proceedings of the National Academy of Sciences* 118, e2105646118. doi:10.1073/pnas.2105646118
- Seth, A. K. (2013). Interoceptive inference, emotion, and the embodied self. *Trends in Cognitive Sciences* 17, 565–573. doi:10.1016/j.tics.2013.09.007
- Seth, A. K. and Bayne, T. (2022). Theories of consciousness. *Nature Reviews Neuroscience* 23, 439–452. doi:10.1038/s41583-022-00587-4
- Stettler, D. D., Yamahachi, H., Li, W., Denk, W., and Gilbert, C. D. (2006). Axons and synaptic boutons are highly dynamic in adult visual cortex. *Neuron* 49, 877–887. doi:10.1016/j.neuron.2006.02.018
- Tang, H., Schrimpf, M., Lotter, W., et al. (2018). Recurrent computations for visual pattern completion. *Proceedings of the National Academy of Sciences* 115, 8835–8840. doi:10.1073/pnas.1719397115
- Tolman, E. C. (1948). Cognitive maps in rats and men. *Psychological Review* 55, 189–208. doi:10.1037/h0061626
- Tononi, G., Boly, M., Massimini, M., and Koch, C. (2016). Integrated information theory: From consciousness to its physical substrate. *Nature Reviews Neuroscience* 17, 450–461. doi:10.1038/nrn.2016.44
- Uehara, M. A., Jacobson, N., and Moussavi, Z. (2024). How accurate are coordinate systems being used for transcranial magnetic stimulation? *Frontiers in Human Neuroscience* 18, 1342410. doi:10.3389/fnhum.2024.1342410
- [Dataset] U.S. Food and Drug Administration (2006). Significant risk and nonsignificant risk medical device studies: Guidance for IRBs, clinical investigators, and sponsors. Final guidance; docket FDA-2006-D-0031

- [Dataset] Valdez, J. (2026a). Canvas engineering. Software and research documentation, CommandAGI
- [Dataset] Valdez, J. (2026b). The shape of experience. Online conceptual manuscript
- van Dyck, L. E., Kwitt, R., Denzler, S. J., and Gruber, W. R. (2021). Comparing object recognition in humans and deep convolutional neural networks—an eye tracking study. *Frontiers in Neuroscience* 15, 750639. doi:10.3389/fnins.2021.750639
- Venkadesh, S. and Van Horn, J. D. (2021). Integrative models of brain structure and dynamics: Concepts, challenges, and methods. *Frontiers in Neuroscience* 15, 752332. doi:10.3389/fnins.2021.752332
- von Holst, E. and Mittelstaedt, H. (1950). Das reafferenzprinzip: Wechselwirkungen zwischen zentralnervensystem und peripherie. *Naturwissenschaften* 37, 464–476. doi:10.1007/BF00622503
- Wainberg, M., Forde, N. J., Mansour, S., et al. (2024). Genetic architecture of the structural connectome. *Nature Communications* 15, 1962. doi:10.1038/s41467-024-46023-2
- Wang, C., Zhang, T., Chen, X., He, S., Li, S., and Wu, S. (2023). BrainPy, a flexible, integrative, efficient, and extensible framework for general-purpose brain dynamics programming. *eLife* 12, e86365. doi:10.7554/eLife.86365
- Wang, Y., Wang, Y., Zhang, X., Du, J., Zhang, T., and Xu, B. (2024). Brain topology improved spiking neural network for efficient reinforcement learning of continuous control. *Frontiers in Neuroscience* 18, 1325062. doi:10.3389/fnins.2024.1325062
- Yamins, D. L. K., Hong, H., Cadieu, C. F., Solomon, E. A., Seibert, D., and DiCarlo, J. J. (2014). Performance-optimized hierarchical models predict neural responses in higher visual cortex. *Proceedings of the National Academy of Sciences* 111, 8619–8624. doi:10.1073/pnas.1403112111