

Icon Dynamics: The Substrate of Iconism

Working document, Iconism Theory project. The substrate thread — the ground on which `icon\_and\_world.md`, `consciousness\_awareness\_spirit.md`, `centre\_of\_awareness.md`, `conceivability.md`, and `ai\_and\_sentience.md` rest. Establishes the framing claim that Iconism is Virtualism applied to the brain-icon case, the three-layered emergence within which icons obtain, the high-dimensional star geometry of icons themselves, the handshake mechanism by which icons form and recall succeeds, and the wave-function-constraint account of downward causation. Source citations are by diary number; the underlying extracts live in the diary files in the Iconism Theory project, navigated by `iconism\_file\_inventory\_v006.md`. Cross-references to Tier 1 documents are by section number in those documents.

__Version 002.__ Supersedes v001. This version folds in the ratified refinement findings F1–F23 (notes v002–v004), the O10 resolution D10 (note v005), and the O5 resolution D5 with F28–F29 (note v006), in a single coordinated pass. A change-log mapping each edit to its driving finding is appended as §12.

0. Why this document exists

A first attempt at the Iconism working-document chain would have led with `icon\_and\_world.md` — the central claim that consciousness

is the indiscernibility of icon and world in the Leibniz sense.

That order, on the model of the Tier 1 mistake that led with paradox before establishing foundations, would have been wrong. Before any work can be done on what icons are and how they relate to the world, an account is needed of where icons sit in the brain-mind story, by what mechanism they form and re-form, how they interact with one another, how they are caused by the brain's activity and how they in turn act on it. Without those mechanics in place, the icon-and-world relation would be suspended over nothing.

This document supplies them. It is the substrate document for Iconism on the same reasoning that put `foundations.md` first in Tier 1: most of what follows uses it. The five working documents that succeed it draw on its apparatus as a matter of course; their sections refer back to its mechanisms rather than redeveloping them.

The document is also where two structural claims that the corpus has been building toward for nearly a decade are stated for the first time in worked form.

The first is the framing claim: that Iconism is Virtualism applied to the brain-icon case, using exactly the same machinery, in many more dimensions, with both attraction (from sameness) and repulsion (from difference) operating in each dimension. Gravity is the degenerate case of this machinery operating in spatial

dimensions only and with attraction only. The qualitative differences between consciousness and gravity — the felt presence, the qualia, the integrated awareness, the felt persistence — all follow from a dimensional jump and an asymmetry of attraction and repulsion that the higher dimensions support. The addendum to Diary 064B, 26 May 2026, states this as an identity rather than an analogy: *the brain's reasoning is the inside-the-brain gravity process, exactly the same process as physical gravity but in unlimited dimensions*. The substrate document is the place where that identity is unpacked.

The second is that recall, deduction, the handshake/socket/parts-supplier account of how mind meets brain, the wave-function-guess-and-confirm mechanism, and downward causation via wave-function constraint are *the same mechanism viewed at different magnifications*. The corpus has the strands; this document fuses them. The fusion is what makes Iconism a structural theory of cognition rather than a collection of useful analogies, and what allows the four working documents that follow to draw on a single underlying account rather than to develop separate ones.

The document closes with the framework's response to Princess Elizabeth's question to Descartes — the interaction problem laid for good. That closing is not a flourish. It is the test that the substrate has been built well enough: if Princess Elizabeth is not answered cleanly, the substrate is not yet right.

1. The framing claim

Iconism is Virtualism applied to the brain-icon case. The same machinery — relationships of sameness and difference between virtual wholes, four-state determination, paradox-driven emergence, virtual facts about systems — operates in both. The framework's account of consciousness is not a parallel framework to its account of gravity; it is the **same framework**, applied to a structurally richer case.

This claim is sharper than the analogy-by-centre-of-gravity that the corpus has been developing since Diary 042 (5 October 2020), where the icon was first compared to a centre of gravity. The analogy has now been re-read, in light of the full Virtualism apparatus, as an identity. The diarist's own statement, from the 26 May 2026 addendum to Diary 064B: **the brain's reasoning is the inside-the-brain gravity process, exactly the same process as physical gravity but in unlimited dimensions.**

1.1 Why this is identity, not analogy

The temptation, on first reading, is to take the claim as a vivid metaphor: consciousness is **like** gravity, in roughly the way a pendulum is **like** a tuning fork. But that reading underrates the structural commitments shared between gravity and consciousness on the Virtualist account. Gravity is, in Virtualism, an emergent fact about a system of masses: the system has a virtual centre,

that centre is a virtual fact (no more than its parts, but genuinely there as the focus of the system), the centre influences the parts through the constraint of being-the-centre-of-this-configuration. The icon-and-awareness machinery has exactly this structure: a system of icon-parts has a virtual centre (awareness), that centre is a virtual fact, the centre influences the parts through the constraint of being-the-centre-of-this-icon-collection.

What differs is not the form of the relation. What differs is the *number* and *kind* of dimensions in which the relation operates, and the *asymmetry* of attraction and repulsion that the higher dimensions support. Gravity is the case where one dimension-type (space) is the only dimension, and one relationship-type (attraction-by-sameness, in the form of having-mass) is the only relationship operating. The icon case is the case where many dimension-types operate together (the spatial dimensions of the object iconised, but also its moral dimensions, its mathematical dimensions, its musical dimensions, its dimensions of wholeness, its dimensions of relatedness to other objects), and where in each dimension both sameness-attraction and difference-repulsion can operate at once. This is the dimensional jump and the asymmetry that, jointly, account for why something with the structural form of a gravitational system also has the felt form of a conscious experience.

The identity is not metaphor because gravity's machinery is not external to consciousness's machinery. Gravity is what the same

machinery does *when restricted to spatial dimensions and to attraction only*. Consciousness is what the same machinery does *when permitted the full dimensional range and the full attraction/repulsion structure*. The two phenomena are not analogous because they happen to share form; they share form because they are operations of the same apparatus under different parameter-settings.

1.2 What the framing claim buys

Three things, each of which the corpus has been edging toward and which this document brings into single statement.

First, it answers the question that Diary 049 (11 August 2022) posed when it introduced the term "Iconic Gravity": is the gravity-like behaviour of icons a metaphor or a literal claim? The framework's mature answer, after another four years of work, is *literal*. The brain is doing gravity, but in many more dimensions than the three spatial dimensions to which physical gravity is restricted. The wheel-and-orbit illustration of §2.4 below makes the same point through a different image: a planetary orbit and a solid wheel are both gravitating systems, but the wheel has rigidity-between-parts that the orbit does not, and that rigidity supports emergent rotational dynamics that the orbit can only approximate. The brain-with-connectome-and-firing-patterns is the dimensional analogue of the wheel: rigid relationships between firing-neurons that support emergent wholes the orbit-analogue of stars-in-a-cluster cannot.

Second, it resolves the Hard Problem in the form the corpus has been pressing it. The Hard Problem in its standard formulation asks: how does subjective experience arise from objective brain-activity? The framework's answer is that the question presupposes a discontinuity that does not in fact obtain. Subjective experience is the same kind of fact as the centre-of-gravity of a system of masses — a virtual fact about a system — only operating in dimensions richer than space and with relationships richer than mass-attraction. The Hard Problem is hard not because consciousness is unlike the rest of the world but because the rest of the world, on the standard physicalist reading, is taken to be flat and relation-free, and consciousness obviously is not. Virtualism's gravity is not flat and relation-free either; it is a system of virtual wholes with centres that are virtual facts. The discontinuity dissolves once Virtualism's account of the rest of the world is on the table.

Third, it gives Iconism a structural commitment that survives even where its phenomenological commitments are contested. A critic may dispute whether the icon really is indiscernible from its worldly object in the Leibniz sense; a critic may dispute whether the centre of awareness really has the felt unity the framework attributes to it; a critic may dispute the specific emergence-claim that parts-in-the-right-relation must produce a whole. But the structural commitment — that consciousness is governed by the same apparatus that governs gravity, the same

four-state determination, the same paradox-and-hypodox dual engine, the same virtual-facts-about-systems ontology — is a single claim that stands or falls on the success of the apparatus as a whole. If gravity comes out right on the Virtualist account, much of consciousness comes out right with it.

1.3 What the framing claim does not buy

It does not buy the actual neural mechanism by which any specific recall happens. That is neuroscience. The framework's claims about icon dynamics are structural: they describe the shape of the mechanism (sockets, parts-suppliers, handshakes, four-state configurations, centring) without specifying which neurons, which patterns, which feedback loops the brain uses to instantiate any particular handshake. The brain does its own work at the layer of firing patterns; the framework's claims are at the layer of the patterns' wholes, which is a different layer from the patterns themselves.

It does not buy the empirical proof that any particular system satisfies the structural conditions for icon-bearing. Whether a particular biological brain at a particular moment is in fact producing icons in the framework's sense, whether a particular non-biological system might in principle be so doing, whether specific clinical conditions involve specific kinds of icon-malformation — these are open empirical questions that the framework constrains but does not settle. Iconism is a philosophical account; it inherits from Virtualism the discipline

of stating its claims as structural and leaving their instantiation to the relevant empirical disciplines.

It does not, in particular, buy the claim that current AI is or is not conscious. The framework permits both readings on different substrate-level assumptions; the working document `ai\_and\_sentence.md` develops the structural argument; the final empirical question is left honestly open.

1.4 The plan of this document

The framing claim, made in this section, is given operational form in the seven sections that follow. Section 2 sets out the layered-emergence picture within which icons obtain: three real emergent levels (material brain, firing patterns, icons-and-mind), each related to the next by Virtualism's emergence apparatus. Section 3 develops the geometry of icons themselves: high-dimensional stars, attraction-by-sameness and repulsion-by-difference, the NATO-star structure that allows fine selectivity in which dimensions sameness holds. Section 4 develops the handshake mechanism in four-state-determination terms, showing how the mind's socket-creation, the brain's activity-layer guessing, and the recentring of awareness on successful icon-formation are the same machinery in four configurations. Section 5 develops downward causation as wave-function constraint, answering Kim's Causal Exclusion challenge and completing the Supervenience sub-section that the 9–14 February 2026 essay-draft (Diary 063B) left as a header. Section 6 collects the worked applications: confirmation

bias, tip-of-the-tongue, recall, deduction, imagination,
perception — all the same mechanism in different configurations.
Section 7 takes a position on the continuity of consciousness
across Arcs, the second hard problem posed in Diary 059A. Section
8 closes with the framework's response to Princess Elizabeth's
question to Descartes: one substance, layered emergence, downward
causation as wave-function constraint — no interaction problem
because there is no second substance to interact across.

2. Layered emergence: brain, activity, mind

Iconism is not a layer of mental stuff sitting on top of brain
stuff. It is a position about what kinds of fact obtain at what
levels of the brain-mind system, and how facts at one level relate
to facts at the level above and the level below. The picture is
of three real emergent layers, each with its own dynamics, each
related to the next by Virtualism's emergence apparatus
(`paradox\_and\_emergence\_v003.md`). The brain is not a thing that
has a mind; the brain-with-its-firing-patterns *is* what minds
do in the dimensions of biological substrate, and minds *are* what
brain-firing-patterns do at the level where indiscernibility-with-
worldly-wholes can hold.

This section sets out the three layers, what is real at each, what
facts each supports, and how the layers relate. The conventions
established here are used by the remaining sections of this
document and by the four working documents that follow.

2.1 The three layers

****Layer 1: The material brain.**** The first layer is strictly material — the neurons, the synapses, the glial cells, the vascular network, the chemical environment, the connectome that links one neuron to another. This is the brain that the neuroscientist studies and the surgeon operates on. Its existence is uncontroversial. Its dynamics — the rules by which neurons fire, the rules by which synapses strengthen and weaken, the rules by which the connectome is laid down in development and modified in life — are the subject matter of neuroscience and the framework does not attempt to specify them.

The material brain's *shape* — its connectome, the wiring pattern that determines which neuron can influence which — is what constrains what activity is possible. Two brains with different connectomes can produce different ranges of activity. The material brain is not itself thinking; it is the substrate that permits thinking-shaped activity to occur.

****Layer 2: The brain's activity.**** The second layer is the patterns of co-ordinated firing that occur within the material brain. A firing pattern is a fact about the material brain — at this moment, this group of neurons is firing in this co-ordinated way — that is not identical with any single neuron's firing or any single synapse's state. The pattern is *of* the firing-neurons, but the pattern itself is not any single firing-neuron. It is a virtual

whole composed of the firings, supervening on the material brain but not reducible to it in any way that the framework cares about.

This is the layer at which neuroscience does most of its interesting work. EEG, fMRI, neural recording, computational modelling — all are studies of layer 2, the patterns of coordinated activity that the material brain supports. The patterns are what the brain **does**; the connectome is what the brain **is**.

The relation between layer 1 and layer 2 is straightforward emergence in Virtualism's weak sense (``paradox_and_emergence_v003.md` §5`): the patterns supervene on the firings, they are no more than the firings collected into wholes, no further fact is required. The pattern is a fact-about-the-firings, not a thing distinct from them. (``foundations_v001.md` §3` establishes that facts can be more than their nodes without being something **over and above** their nodes: the pattern is more than any individual firing, in that it involves the relations between firings, but it is not a further existent on top of the firings. It is what they amount to, collected.)

****Layer 3: Icons and mind.**** The third layer is the wholes formed by the firing patterns at the level where indiscernibility with worldly wholes can hold. This is the layer at which icons obtain. An icon is not a firing pattern; it is the virtual whole that the firing pattern produces at the level of indiscernibility. The firing pattern is what the brain is doing at layer 2; the icon is

what that doing amounts to at layer 3, where the doing has the structural form of a whole indiscernible from a worldly whole.

This is the layer that Iconism is principally about. Consciousness is the obtaining of icons at this layer. Awareness is the centre of consciousness at this layer. The framework's claim is that this layer is real — not in a sense that requires it to be ontologically extra, but in the same sense that the centre of gravity of a gravitational system is real: a virtual fact about the configuration of the layers below, no more than those layers but genuinely there as the focus of them.

Spirit, the third term of the triad that ``consciousness_awareness_spirit.md`` develops in full, is the collection of all true facts about the individual — a fact rather than a reality, and the term that carries the Iconism–Spirit boundary. The rider to record here, where the three terms are first in view, is that mind and spirit are **identical in kind but not in extension** (F18). They are the same kind of item — virtual facts, icon-wholes of the sort this layer obtains — differing only in reach: mind is the live, currently-obtaining icons that the brain is lighting now; spirit is the whole standing collection of an individual's true facts, the lit-now included as its leading edge and the no-longer-lit retained as eternal fact. The ontology of that retained collection is Spirit's to develop; what matters here is only that nothing of a different **kind** is being introduced when the talk turns from mind to spirit.

The relation between layer 2 and layer 3 is more interesting than the relation between layer 1 and layer 2. It is still emergence, but emergence in the framework's **strong** sense: an icon does not merely collect firing-patterns into a whole, it stands in a specific relation (indiscernibility-with-a-worldly-whole, in the Leibniz sense) that the firing-pattern alone does not stand in. The whole-formation here is paradox-driven in the sense of ``paradox_and_emergence_v003.md` §2`: the firing pattern, treated just as a firing pattern, is overdetermined as a description of what the brain is doing — every firing has a specific time, a specific neuron, a specific synaptic context, and the totality of those facts forces a level-change. The level it forces is the icon: the whole that the firing pattern produces at the dimensions where indiscernibility-with-the-worldly-object can hold.

This is precisely the relation between gravity and its parts, applied to a higher-dimensional case. A system of masses has, by Virtualism, a virtual centre — this is forced by exactly the arguments Doc one and Doc two make (``gravity_and_the_past_v1_0.md` §13.6`). A system of firing-neurons has, by exactly the same kind of argument extended to higher dimensions, a virtual whole at the level of icon-relations. The emergence at the brain-to-mind layer-join is the same kind of emergence as the emergence at the parts-to-centre-of-gravity join, only operating in dimensional richness that the gravity

case lacks.

The three layers are the *lower half* of a single constraint ladder, and it is worth running the ladder up to its top before moving on, because the upper rungs are what the rest of this document and the chain that follows are about. Read upward, the ladder is: activity → icons → consciousness → awareness. Each rung is a whole standing to the rung below as the centre of gravity stands to its masses — forced into being wherever the configuration below forces a whole (the forced-emergence point that answers Strawson, F20), and constraining the rung below in turn. Firing patterns, collected at the dimensions where indiscernibility can hold, force icons; icons, obtaining together, are consciousness; consciousness, centred on its most-connected peak, is awareness. The ladder does not climb forever: awareness is its top rung, the pinnacle of the configuration, the point past which there is no further whole to form because there is nothing left for a further whole to be the centre *of*. The downward causation of §5 is just this same ladder read in the other direction — the upper rungs constraining the lower — and the centre-of-awareness treatment of `centre\_of\_awareness.md` is the account of the top rung in particular.

2.2 What is real at each layer

The framework's commitment is that all three layers carry real facts. The material brain has real facts: there are neurons here and not there, this synapse is excitatory and that one inhibitory,

the connectome runs this way and not that. The activity layer carries real facts: this firing pattern occurred at this time, those neurons fired together, this oscillation has this frequency. The icon-and-mind layer carries real facts: this icon stands in indiscernibility-relation with that worldly object, this icon is the centre of awareness, this icon is more strongly connected to those icons than to those others.

The framework does **not** commit to the claim that the layers are ontologically separate. Layer 1 is the material brain, in all its specifics. Layer 2 is what the material brain is doing, considered as patterns of co-ordinated firing. Layer 3 is what those patterns amount to at the level of indiscernibility-with-the-world. They are not three substances; they are three levels of fact about one substrate, related by emergence-relations the framework has characterised.

This matters because it is what makes the layered-emergence picture consistent with Virtualism's one-substance ontology ('foundations_v001.md` §6: everything reduces to virtual facts). There is no second substance for the mind to be made of, in addition to the material substrate. There is one substrate (however its ontological status is itself analysed by Virtualism) and three levels of fact-about-the-substrate, the upper levels emergent from the lower in the framework's structured sense.

The framework also does not commit to the claim that the layers

exhaust what is real about the brain. There may be other layers the framework does not currently develop. There may be sub-layers within each of the three named. The commitment is only that *at least* three layers must be distinguished if the framework is to work, and that these three are sufficient for the substrate-document work.

2.3 What facts each layer supports

This is the discipline that prevents layer-conflation in subsequent sections. Each layer supports facts of a particular kind, and the framework's claims about each layer respect this.

****Layer 1 supports material facts.**** Locations of neurons, strengths of synapses, properties of myelin sheaths, oxygen gradients in the vascular bed, distributions of neurotransmitters. These are facts of biology and biochemistry. They are causally implicated in everything that happens at layers 2 and 3, but they are not themselves at those layers. Neuroscience studies these facts; the framework defers to neuroscience on them.

****Layer 2 supports activity facts.**** Firing patterns, oscillation frequencies, synchronisation events, network states. These are facts about what the brain is doing, considered as a pattern-producing system. They supervene on layer 1 but are not identical with any specific layer-1 fact. Computational and pattern-analytic neuroscience studies these facts. The framework defers to this discipline on them, with one exception: where a layer-2

fact stands in a specific relation to a layer-3 fact (specifically, where a firing pattern produces an icon), the framework has something to say about the relation, even if not about the layer-2 fact itself.

****Layer 3 supports icon facts.**** Icons obtaining or failing to obtain, icons standing in indiscernibility-relations with worldly wholes or failing to do so, icons being more or less central to the awareness-collection, icons being in deduction-relations or recollection-relations or imagination-relations. These are facts of mind. Philosophy studies them. The framework's claims are principally about this layer and about its relation to the layer below.

The icon at layer 3 is not a **kind** of firing pattern at layer 2. The icon is what the firing pattern amounts to at a higher level of organisation. The relation is not classificatory (one pattern counts as one icon, another pattern counts as another) but emergent (the pattern produces the icon at a level where the icon can stand in relations the pattern alone cannot).

2.4 The wheel-versus-orbit illustration

The framework's three-layer picture is unusual enough that an illustration helps. Two systems, both gravitational, both analysable in Virtualism's terms, but supporting different ranges of emergent dynamics.

****The planetary orbit.**** A system of two masses in mutual gravitational interaction. Each mass has its own properties — its mass, its position, its momentum. The system has a centre of gravity, which is a virtual fact about the configuration of the two masses (``gravity_and_the_past_v1_0.md` §13.6`). The orbit is the dynamics that result from the masses' mutual interaction through their respective positions relative to the centre of gravity. The centre of gravity exists, it influences the orbit, but it is no more than the configuration of the masses; if the masses were rearranged, the centre would be elsewhere.

The orbit has limited emergent properties. The system has only spatial dimensions to operate in. Each mass interacts with the other only through the gravitational relation (in the simple case). The system has no rigidity beyond what gravity provides; the masses can drift apart, change their orbital parameters, exchange positions, without the system as such resisting these changes. There is no whole that the masses are bound into; there are just two masses with a centre of gravity that moves with them.

****The solid wheel.**** A second system, also gravitational in the sense that it has mass and so contributes to gravity, but with an additional structural property: rigidity. The parts of the wheel are bound together by chemical bonds, lattice forces, friction between surfaces — by structural relations that link each part to its neighbours and so to all parts. When one part of the wheel

moves, the others move with it; the wheel rotates as a whole.

The wheel has emergent properties the orbit cannot have. It has rotational dynamics — it can store rotational energy, it can transmit torque, it can gear with other wheels in a machine. None of this is available to the orbit, which has none of the binding relations the wheel has. The wheel's emergent dynamics depend on its parts being in **binding** relations with each other, not merely in gravitational ones. The same physics underlies both systems; the wheel just has additional kinds of relation operating among its parts.

****The brain-with-connectome-and-firing-patterns is the wheel-analogue, not the orbit-analogue.**** The neurons of the brain are not free-floating gravitational masses in mutual attraction; they are bound together by the connectome — a network of synaptic connections, supplemented by local chemical environments and oscillatory entrainment, that links each neuron's firing to other neurons' firings in highly specific ways. The firing patterns of the brain are not merely the gravitational analogue of "co-occurring activities"; they are the rotational analogue, the emergent dynamics that the bound system supports. Just as a wheel can rotate where an orbit cannot, a brain can produce icons where a system of unbound activities cannot.

This is the analogical illustration of the layered-emergence picture. The relations among the brain's parts are the binding

relations the wheel has; the emergent dynamics those binding relations support are the icons that the brain produces; the centre of awareness is the rotational analogue of the wheel's centre of rotation, a virtual fact about the system that has the same emergent reality the wheel's rotation does. The Sperry-wheel emergent-rotation analogy (Diary 058, 27 November 2024) makes this point in the corpus's own vocabulary; the wheel-versus-orbit distinction sharpens it by contrast with a case that **lacks** the binding relations.

The illustration is meant to be read structurally rather than substantively. The framework is not claiming that consciousness is literally a kind of rotation. It is claiming that consciousness stands to brain activity as rotation stands to translation in a bound system: emergent dynamics that the binding relations support but that the parts taken individually do not. The same physics, in both cases — gravity and (in the gravitational case literal) attraction; sameness-attraction and difference-repulsion in the brain case — operating on parts that are bound together in ways that support an emergent layer with its own dynamics.

2.5 What the framework claims and does not claim

The discipline of working at this layer of abstraction is to be clear about what the framework is in a position to commit to and what it is not.

****The framework claims:****

- The structural shape of how mind and brain-activity meet (the handshake/socket/parts-supplier account of §4, sharpened in Diary 064B, 12 May 2026).
- The structural reason there is no interaction problem in Princess Elizabeth's sense — one substance (Virtualism's relationships-between-virtual-wholes), layered emergence, downward causation as constraint on which emergent outcomes obtain (§5 and §8).
- The conditions under which an icon can form: indiscernibility with a worldly whole, in the dimensions the icon participates in (developed in `icon\_and\_world.md`, presupposed here).
- The structural verdicts on the five p-cases jointly consistent with the four-state-determination apparatus (developed in `conceivability.md`, presupposed here).

****The framework does not claim:****

- The actual neural mechanism by which any specific recall happens. That is neuroscience. The framework's claims are structural, not biological.
- The specific feedback loops, correlations, oscillation frequencies, network configurations the brain uses. That is neuroscience.
- That every brain produces icons in exactly the same way. Multiple realisation is preserved at layer 3: the same icon can be produced by different firing patterns, because what

individuates the icon is its layer-3 relation (indiscernibility-with-a-worldly-whole), not the layer-2 substrate that supports it. The framework's commitment is stronger than functionalism's because it identifies *what* the common feature is (the icon as virtual whole indiscernible from a worldly whole), not merely that something common exists; but it is at the same time more permissive than any identity theory, because the substrate is not what individuates.

- That the layered emergence picture exhausts the architecture of the brain-mind system. Other layers may exist; the framework names the three it needs.

This discipline is what makes the framework presentable to working philosophers and to working neuroscientists without overreach. The author's recurring honest disclaimer — *I am not a neuroscientist* — is not a weakness but the correct shape of a philosophical account: it stakes its claims at a level of generality the empirical disciplines do not foreclose and leaves the level below to those who are equipped to investigate it. This is what `foundations\_v001.md` §9 calls the *what this document does not do* discipline, applied to consciousness.

3. The icon as high-dimensional star

The icon is a virtual whole at the third emergent layer of the brain-mind story. Its geometry is what makes it the kind of whole

that can stand in indiscernibility-relations with worldly objects in many different dimensions at once. This section develops that geometry: the dimensional jump, the asymmetric structure of attraction-by-sameness and repulsion-by-difference, and the high-dimensional-star form that the geometry takes.

3.1 The dimensional jump

Gravity operates in spatial dimensions only. A system of masses has positions in space; its centre of gravity has a position in space; the gravitational interactions among the masses are mediated by spatial relations. The whole gravitational story is played out in three dimensions (or, in relativistic settings, four — but still a small finite count).

Icons participate in many dimension-types, and in many dimensions within each type. An icon of a pear sitting on a table is not located only in the spatial dimensions of the pear's position; it is also located in the dimensions of the pear's colour (a multi-dimensional space of hue, saturation, brightness, ripeness-appearance), of the pear's mass (a one-dimensional space, but real), of the pear's edibility (a structurally complex space of nutritional, gustatory, and historical-personal dimensions), of the pear's relations to other pears (a space of similarities and differences across the perceiver's experience of pears), of the pear's role in the present scene (Pear-as-object-of-attention, Pear-as-still-life-element, Pear-as-tomorrow's-breakfast). Each of these is a real dimension along which the icon has structure.

The icon of a complex object — a person, a building, an argument, a piece of music — has dimensions correspondingly more numerous.

The icon of an argument has dimensions of logical structure, rhetorical structure, the persons engaged, the history of the disagreement, the stakes of the resolution. The icon of a piece of music has dimensions of pitch, rhythm, timbre, harmony, mood, narrative arc, the performer's character, the historical moment of composition. There is no fixed upper bound on the dimensional count.

This is the *dimensional jump* the framing claim refers to. The same machinery that produces a centre of gravity for a system of masses in three spatial dimensions produces a centre of awareness for a system of icons in a much larger number of dimensions of many different kinds. Centre-of-gravity machinery in three dimensions is gravity; centre-of-icons-machinery in many dimensions is consciousness. The two phenomena look unlike each other only because their dimensional richness differs. The underlying operation is the same.

The Log v6 statement is precise: icons *are abstract in the strict sense — they have no specific position in the brain because their correlates are dispersed throughout it — and they are multidimensional, having as many dimensions as the concepts they instantiate*. The framework develops this beyond the Log: the dimensional count is not just "many" but *as many as the concepts

the icon instantiates*, where the concepts include not only the properties of the worldly object but the relations of that object to other objects, to the perceiver, to the perceiver's history, to the situation at hand. Icons inherit the full dimensional structure of what they are iconic of.

The diarist's 30 April 2026 speculation on mental-particles (Diary 064A) — abstract carriers of "mental force" paralleling photons and gluons but in higher dimensions — is the first attempt in the corpus at a *mechanism* for how icons interact across these dimensions. The substrate document does not commit to the mechanism but flags it as candidate substrate for future work: see §7.5 below. What the document does commit to is the dimensional structure itself, which is independent of any particular force-carrier picture.

****What every icon-dimension is, and what the icon instantiates.****

Before the dynamics, one structural point must be fixed, because the working documents and the core claim both rest on it. *Every dimension of an icon is a pointer; the icon instantiates none of them.* The icon of the pear does not contain a pear's mass, a pear's greenness, or a pear's position; it carries, in each dimension, a generated handle that points at the corresponding feature of the world. The mass-dimension of the icon is not a small mass in the head — it is a created pointer toward the pear's real mass (F1). The same holds, dimension by dimension, across the whole icon.

What then makes some dimensions **shared** and others **created**?

Every icon-dimension alike is a fact generated, at the third layer, from the same incoming data — the same photons that deliver colour deliver shape and position too; none is a direct, unmediated read-off of the world (F28). The difference is whether the generated pointer points at a Tier-1 fact **of its own kind**, one it can be indiscernible from:

- **Position and shape are shared, and they anchor.** There is a real Euclidean geometry out there, with prior existence; the rendered icon-fact is indiscernible from it by Leibniz's Law, and that indiscernibility is what ties the icon to a real object. Shape needs one further word, because it looks as though it sits on both sides at once. It does not, once the **whole** is distinguished from the **parts-and-relations** (D5): the metric geometry — the parts and their relations — is the shared, anchoring fact; the unified **seen** shape — the rendered gestalt, the whole — is the mind's creation. Shape exists as an abstract in Tier 1 but must be **rendered** in Tier 2 to be known as an icon, whereupon it is immediately recognisable by Leibniz's Law.
- **Colour is created, and it is added.** There is a Tier-1 optical property — which wavelengths the object sends to the eye — but the brain constructs colour from the proportions in which its three (sometimes four) cone-types are triggered, much as a screen constructs colour from RGB values. The resulting quale has no Tier-1 correlate of its own kind: non-spectral hues such as

magenta have no wavelength at all, so colour cannot be in the light. Colour is therefore a Tier-2-only fact, created — a true fact the mind **adds** to Existence (F2), true-but-not-real, a handle onto a real property we could not otherwise register. Because that added fact is nonetheless a **standing** Tier-2 fact, it is available to be remembered, which is what the account of memory below requires.

The icon thus **spans** the determination scale rather than sitting at one point on it, and it does so by traversal, not by static occupancy (D10). Its anchoring runs through the shared dimensions, which are over-determined — the present, solid, boosted object. Its creation is the outward extension: the created dimensions are under-determined pointers, addable and removable, instantiating nothing, reaching toward the non-determined void — the real properties we cannot detect, faintly indicated and never filled. An icon forms by moving through that scale: an under-determined rearrangement at the activity layer, a closure into one whole (the configuration forcing a single icon), and the lit settled fact of the perception. §4.3 sets the traversal out as motion along the scale; here the point is only that "shared" and "created" name the two ends the icon reaches between, not two separate kinds of icon.

This fixes the core claim of the chain, which ``icon_and_world.md`` carries in full and which this document states once, here, as the ground of everything downstream:

- > Consciousness is the obtaining of an icon that is indiscernible
- > from its worldly object on the dimensions they share — geometry,
- > position, relation, by Leibniz's Law — and that exceeds it on the
- > dimensions consciousness creates: colour, felt-weight, the sensory
- > qualities. The shared dimensions *anchor* the icon to a real
- > object; the created dimensions are what consciousness *adds* to
- > Existence. Indiscernibility is therefore not the whole of
- > consciousness but its anchor: it is what makes the icon an icon
- > *of something* rather than a free-floating invention, while the
- > created qualities are what make it an experience.

A gloss the claim needs, lest "worldly object" be read too narrowly: the world-anchored perceptual case is *primal* — the evolutionary origin of icons (F22). Icons extend from it to the self, the past, and invented or ideal worlds; along that extension the anchor weakens and the created component dominates, with fiction the limit case — a partially untrue fact with no worldly object to be indiscernible from. "Worldly" names the primal anchored case, not a restriction of icons to external physical objects.

3.2 Sameness-attraction and difference-repulsion

In each dimension an icon participates in, the icon stands in relationships with other icons. Two icons of pears stand in relationships across all the dimensions both icons participate in: their colour-dimensions, their position-dimensions, their weight-dimensions, their context-dimensions. Within each dimension, the relation between the two icons can be one of

sameness or one of *difference*.

`foundations\_v001.md` §5 establishes Leibniz's Law as the law of virtual connection: where two objects are indiscernible in some feature, that feature is — by the Law — *the same feature*, not two similar features. Two red balls share *one* redness, not two similar rednesses. This is the foundation-level account of sameness, applied to features-of-icons in the present case. Where two icons are indiscernible in some feature, that feature is the same feature in both icons.

Indiscernibility creates a relationship the framework can call *sameness-attraction*. Two icons sharing a feature are bound together at that feature; they form a sub-whole at the dimension in which they share. This is the foundation-level mechanism of relationship-formation

(`paradox\_and\_emergence\_v003.md` §6, *Stage 1 — Sameness*) applied to icons.

Distinctly, where two icons differ in some feature, the feature distinguishes them. They are *not* the same icon at that feature; the difference is what makes them two icons rather than one. This is a relationship in its own right, the framework's *difference-repulsion*: two icons differing in some feature push apart at that feature, keeping the icons distinct.

Gravity is the degenerate case in which the only relationship-

type operating is sameness-attraction (in the form of mass-attraction: each mass is "the same" as each other mass in being massive, and so they attract). Mass-difference does not produce mass-repulsion in the gravitational case; that asymmetry is what makes gravity always attractive. The icon case is the case where both relationship-types operate together in each dimension. An icon of a green pear and an icon of a green apple are bound together by sameness at the colour-green dimension; they are pushed apart by difference at the pear/apple shape-dimension. The two icons are simultaneously attractive at some dimensions and repulsive at others, the net relation depending on the configuration of sameness-features and difference-features between them.

This is what the framing claim's *asymmetric attraction and repulsion* amounts to. Gravity, with attraction only, is the simplest case. Consciousness, with attraction-by-sameness and repulsion-by-difference operating together in each of many dimensions, is the richest case. The richness is not a different mechanism; it is the same mechanism with the full attraction/repulsion structure operating in many dimensions.

3.3 The geometry: high-dimensional stars

Icons participate in many dimensions and stand in relations of sameness or difference in each. What shape does this give them?

The answer is *high-dimensional star*. Not a sphere — a sphere

would be the shape of an icon whose substance was distributed evenly across all its dimensions, which is not what icons are like. The icon of a pear is not equally about pears, about graspability, about the colour green, and about being on a table; it is **most strongly** about the specific configuration that this particular pear instantiates — its specific green, its specific position, its specific relation to the table — and less strongly about the larger spaces those specific features inhabit. The icon has **points**: specific configurations along specific dimensions where its content is concentrated.

A star, in this image, is a high-surface low-volume shape that has its substance close to its boundary, with points reaching out into particular dimensions. In two dimensions, a five-pointed star has five points reaching outward and a central body that is smaller in area than the convex hull of the points would suggest. In higher dimensions, the geometry becomes more extreme: a high-dimensional star has very little volume relative to its surface area, and its points reach into specific configurations in specific dimensions while its "body" is much closer to the boundary than to a central interior.

The NATO-star image makes this vivid because the visual familiarity of the NATO compass-rose insignia carries the geometric idea forward: substance concentrated near the points, each point in a particular direction (dimension), much of the star's volume close to where the points reach. The icon has this

shape because its content is **specific** in each dimension — it is the icon of **this** pear, with **this** green, on **this** table — and the specificity concentrates the icon's substance near the particular configuration in each dimension that the worldly object instantiates. Each point is a **pointer** at that configuration, not a copy of it (§3.1): the weight-point is a created handle aimed at the pear's mass, the green-point a created handle aimed at the wavelengths it sends — the icon instantiates neither, it reaches at both.

The geometry has three immediate consequences.

****First, icons are highly sensitive to small differences near their points.**** A small change in the specific green of the pear is a change near a point of the icon, and this changes the icon substantially — the icon becomes the icon of a slightly different pear, in a way that matters more than the corresponding small change in some dimension the icon is not pointed in. Icons are not robust to all small changes equally; they are robust to small changes far from their points (the pear could be on a slightly different table-position without much change to its icon) and sensitive to small changes near them.

This is what gives the icon its **fineness** of content. The icon is not a fuzzy schematic of pears in general; it is a sharp representation of **this** pear, with structure concentrated in the particular configurations the pear instantiates. The brain's

work, at layer 2, is to produce firing patterns whose layer-3 icons have this fineness.

****Second, indiscernibility between icons holds at the points.****

Two icons that share a point — two icons of pears that share the specific green, say — are indiscernible at that point and so are, by Leibniz, the same at that point. The sameness-attraction operates at the points where indiscernibility holds; the difference-repulsion operates at points where they reach in different directions. Two icons can be indiscernible at some points and discernible at others, and the icon-icon relation has the configuration of these point-by-point matches and mismatches.

This is what allows the icon of **this** pear and the icon of **that** pear, which are clearly different icons, to nevertheless share content with each other — they may share the pear-shape point, the pear-edibility point, the typical-pear-colour point, while differing at the specific-this-pear-position point and the specific-this-pear-ripeness point. The two icons stand in a rich configuration of sameness-and-difference relations across many points, which is what makes them recognisably icons of the same **kind** of object while remaining icons of different particular objects.

****Third, the centre of an icon-collection is concentrated near the points most-shared.**** Where many icons share a point, the sameness-attraction at that point is strong. The centre of the

collection — the awareness-centre, in ``centre_of_awareness.md``'s terms — is biased toward the configurations of points where sameness is dense. In ordinary perception, where the icons in play are all icons of objects in the present scene, the centre sits near the configurations of points the present scene concentrates: this scene, here, now, with these objects in these relations. The centre is not in a single icon but in the configuration of where the icons' points cluster.

3.4 Why indiscernibility is selective

The high-dimensional-star geometry explains a feature of indiscernibility that the straight Leibnizian reading might otherwise obscure. Indiscernibility between two icons does not have to be total. Two icons can be indiscernible at one point and discernible at all others, and Leibniz's Law applies at the point where they are indiscernible: the feature there is the same feature in both icons.

This is the foundations-level claim of ``foundations_v001.md`` §5 — that Leibniz's Law applies to features rather than only to whole objects — applied to the icon case. The icon of a pear and the icon of an apple share the green-colour point (in the case of a green apple) and so are the same feature at that point; the framework's commitment is that this is not two similar greens but one green, instantiated in two icons. The icons remain two icons because they differ at most other points; the **green** in them is one green because they are indiscernible

there.

This selectivity is what allows the icon to be indiscernible
from its worldly object in the way Iconism requires. The icon
is not in toto indiscernible from the pear in toto — the icon is
in a brain, the pear is in the world, they differ in their
substrates, their causal histories, the ways they can be acted on.
The icon is *indiscernible at the points where icons can be
indiscernible from worldly wholes* — at the configurations of
form, content, relational position, that the icon shares with the
pear and that constitute what the icon is iconic of. The points
at which the icon and the pear are indiscernible are the points
the icon is *about*; the points at which they differ (substrate,
causal access) are not what the icon is about.

The selectivity is what makes Iconism a workable account. If
indiscernibility had to be total, no icon could ever be
indiscernible from any worldly object — every icon would differ
from its worldly object in substrate. The selectivity, grounded
in the geometry, gives the indiscernibility a structural shape
that is achievable for actual brains.

`icon\_and\_world.md` develops the indiscernibility relation in
detail; the substrate document needs only to establish that it is
selective, point-based, and grounded in the geometry of icons as
high-dimensional stars. That establishment is what this
subsection has done.

3.5 Three-photons-three-neurons: a worked geometric example

Diary 059A (11 January 2025) gives the cleanest worked example in the corpus. Three photons of three different colours strike three different neurons; each neuron fires in response to its photon. The firing pattern is at layer 2 — three neurons in co-ordinated firing. The icon of the triangle (the configuration the three photon-sources are in, considered as a whole) is at layer 3. The question the diarist poses, and answers, is: where is the triangle?

The triangle is not in any single neuron — each neuron knows only about its own photon. The triangle is not in any single photon — each photon carries only its own colour and direction. The triangle is not in the firing pattern as such — the firing pattern is three firings in temporal proximity, with no internal feature that announces "and these firings together constitute a triangle". The triangle is at layer 3, in the icon that the firing pattern produces: the virtual whole whose virtual wholeness is indiscernible from the virtual wholeness of the configuration the three photon-sources are actually in.

The diarist's analogy is to the centre of gravity. The icon of the triangle, like the centre of gravity of three masses, is a virtual fact about the configuration of layer-2 parts. Just as the centre of gravity has no position of its own (it is exactly where the configuration of masses makes it be, and only there), the icon of the triangle has no position of its own (it is

exactly the triangle the photons make, and only that). The iconism extension to this analogy, which the substrate document makes explicit, is: the centre-of-gravity machinery and the icon-of-the-triangle machinery are *the same machinery*, only operating in different dimensional ranges. The three masses are in three spatial dimensions; the three photon-sources are in three spatial dimensions *plus* the dimensions of colour, of triangle-shape, of being-a-triangle-rather-than-some-other-configuration. The dimensional richness is greater; the operation is the same.

The worked example also shows the high-dimensional-star geometry in action. The icon of the triangle is concentrated at a particular configuration of points: this specific triangle, with these specific colours at its vertices, in this specific configuration relative to the observer. Small changes near the points (a slight shift in the colour of one vertex, a slight change in the angle at one vertex) change the icon substantially; small changes far from the points (a slight change in the overall position of the triangle in space, if that is not what the icon is about) leave the icon largely unchanged. The icon's substance is concentrated near its points, which is what makes it the icon of *this* triangle rather than a fuzzy schematic of triangles in general.

3.6 What this section has established

Icons are virtual wholes at the third emergent layer of the

brain-mind story. They participate in many dimensions of many kinds — spatial, colour, relational, contextual, conceptual.

Within each dimension they stand in relations of sameness-attraction or difference-repulsion with other icons, and the indiscernibility-relations are at specific points where icons share content. The geometry of icons is high-dimensional star: substance concentrated near points, where each point reaches into a particular configuration along a particular dimension.

This geometric account is what `icon\_and\_world.md` will develop into the full theory of icon-and-world indiscernibility. The substrate document needs only the geometric outline, which is what this section has supplied. The geometry is what makes icons fine-grained, what makes their indiscernibility-with-the-world selective, and what makes the configuration of icon-icon relations rich enough to support the awareness-centring of the next section.

4. The handshake mechanism

The icon's geometry — high-dimensional, point-based, indiscernible with the world at the points where it is iconic — is the static picture. The icon, as a structure, has the form described in §3. But icons are not always already in place. They form, fail to form, re-form, dissolve, and re-emerge in the ordinary course of cognitive life. A workable substrate document needs to say **how**.

This section gives the mechanism. The account in the corpus crystallises around the 12 May 2026 **consciousness-as-handshake** passage of Diary 064B, which describes the mind as creating a **socket** and the brain as supplying **parts** to fill it. The 17 January 2025 **Memory** passage of Diary 059A gives the two-stage **wave-function guess-and-confirm** account that pairs with this: the brain proposes a candidate, the prior memory's wave-function-contribution carries the proposal to a fit. The substrate document fuses these accounts and shows that the resulting mechanism is precisely the four-state-determination apparatus of ``paradox_and_emergence_v003.md` §3`, applied to the mind-brain relation.

The unifying claim of the section: recall, deduction, perception, and imagination are **the same mechanism** viewed under different conditions. The mechanism itself is the handshake. The conditions that distinguish recall from perception from imagination are conditions on what shapes the socket.

4.1 The mind layer creates a socket

The first move of the mechanism, at any moment when the mind is about to do something new — recall, perceive, deduce, imagine — is for the mind layer to **create a socket**. A socket is an under-determined whole at layer 3: it has the structural shape that an icon would have to satisfy for the recall (or perception, or deduction, or imagination) to succeed, but it does not yet have the content that would constitute such an icon.

This is the icon-shaped *hole* in Iconist vocabulary, or the *abstract hole left in the questioning mind by past truth* (Log v6 *p-psychic* passage). The mind has the shape of the answer before the answer is in place; the work of cognition is to fill the shape with the matching content.

The socket has dimensional content — it specifies what kind of icon would fit, in which dimensions, with what kind of points. In recall of a friend's name, the socket has dimensions of name-shape (length, rhythm, initial consonant if remembered), of referential content (the person being recalled, with whatever of their characteristics the mind currently has access to), of phonological-rhyme-with-other-known-names, and so on. In perception of a scene, the socket has dimensions of present-moment-spatial-arrangement, of expected-types-of-object, of recently-attended-features. In deduction, the socket has dimensions of logical-structure-the-conclusion-must-satisfy, of relational-form-given-the-premises, of consistency-with-the-already-believed. In imagination, the socket has dimensions less externally constrained — the dimensional shape is freer, allowing the mind to specify the socket in directions not currently constrained by sensory input or by deductive requirement.

The socket is *under-determined* in the framework's technical sense (``paradox_and_emergence_v003.md` §3, *Hypodox: the underdetermined case*`). It specifies enough that some candidate

icons would fit and others would not; it does not specify enough that exactly one icon must fit. There is room. The brain's task is to find an icon that fills the room.

The shape of the socket is itself a function of what the mind is currently about. In recall, the socket is shaped by the cue (the question, the prompt, the partial fragment that initiated the recall attempt) plus the spirit-facts whose wave-function-contribution is active in the brain's vicinity (the remembered-but-not-currently-conscious facts about the person to be recalled, whose facts are part of what gives the cue specificity). In perception, the socket is shaped by the current sensory input (the photons hitting the retina, the vibrations reaching the cochlea) plus the icons currently in the awareness-collection (which provide the contextual framing in which the new input is interpreted). In deduction, the socket is shaped by the icons already in play (the premises, the goal of the deduction) plus the structural rules of the inference being attempted. In imagination, the socket is shaped by the imaginer's choices about what to specify and what to leave open.

The socket is a layer-3 phenomenon. It is a fact at the icon layer about what icon would satisfy the current task. It is not itself an icon — it has no content, it lacks the point-structure that an icon must have to be the icon of something. It is a specification, an absence-shaped-like-the-answer.

4.2 The brain layer does its own thing

While the mind layer is holding the socket, the brain layer is doing its own work. Neurons are firing, synapses are activating, patterns of co-ordinated activity are forming and dissolving on timescales the framework does not specify in detail. The patterns the brain produces are driven by the brain's own dynamics — by its connectome, by its current chemical state, by the ongoing sensory input, by the patterns of activity that have most recently occurred. The brain is not, in any direct sense, "trying" to fill the socket. The brain is doing what brains do; the firing patterns arise from the brain's mechanics.

The framework does not claim to specify these mechanics. They are the subject matter of neuroscience: the rules by which firing spreads through the connectome, the rules by which oscillatory synchronisations form and break, the rules by which one pattern gives way to another. Whatever those rules are, they are what produces the candidate patterns from which the icon will (if all goes well) emerge.

What the framework does claim is that the patterns the brain produces have **layer-3 emergent consequences**. Every co-ordinated firing pattern produces an icon at layer 3 — a virtual whole at the dimensions where indiscernibility-with-worldly-wholes can hold. Most of the icons so produced are short-lived, weakly structured, not particularly close to what any current socket

calls for. But occasionally a firing pattern produces an icon whose structure happens to fit the shape of the currently held socket. When that happens, the handshake completes — but only when the wave-function-constraint mechanism of §5 has made fitting candidates more probable than the alternatives.

The brain's role in the handshake is to be the **parts-supplier**.

The mind has the shape; the brain has the parts. The brain's parts are firing-patterns; the firing-patterns produce icons at layer 3; the mind's socket is at layer 3; the question is whether the parts produce an icon shaped like the socket.

This is the structural shape of the 12 May 2026 **consciousness-as-handshake** passage (Diary 064B), in the corpus's own vocabulary: **mind as socket, brain as parts-supplier**. The substrate document inherits the metaphor and refits it into the four-state-determination apparatus.

4.3 The four states of the handshake

The four states of determination from ``paradox_and_emergence_v003.md`` §3 describe the configurations of the handshake directly. Each state is one of the four configurations in which a socket-and-candidate-icon relation can sit.

Two readings of the four states are needed, and they are complementary. Read **diachronically**, the states are ***motion along*

a scale^{**}: a single handshake passes from not-knowing, through an open socket, through the press of too-many or too-few candidates, to the settled fit — the determination-status of the configuration changing as the handshake proceeds. Read **synchronically**, the same four states are the standing ***anatomy*** of the scale itself, all regions present at once: a recall situation, frozen at an instant, has its settled fact (the past it is reaching back into), its solid present (the socket as it now stands), its open field of candidates, and its undelineated floor of what has not been called for at all. In this second sense a recall situation is a little Mandala — the whole scale standing still. The subsections below describe the states as configurations; the thing to carry forward is that they are stations on one scale, not four separate boxes, and that an icon-formation is a passage along it.

Non-determined: not-knowing. No socket has been created. The mind layer is not engaged with any particular task; no question is being held; no recall is being attempted; no perception is being framed; no deduction is in progress; no imagination is under construction. This is the structural floor of the handshake mechanism: the minimum condition under which the mind is not currently asking the brain to do anything in particular. The brain's activity continues — neurons fire, patterns form — but none of those patterns is being measured against a current specification.

In ``paradox_and_emergence_v003.md``'s terms, this is the **bare*

co-existence* case: the brain and mind are both in operation, but there is no relationship between them at this level — no socket, no candidate, no handshake under way.

****Under-determined: socket open, candidate not yet emerged.**** A socket has been created. The mind layer is asking the brain for an icon of a particular shape. The brain layer's firing patterns are continuing, and from time to time produce icons that might fit the socket — but no icon has yet emerged whose structure matches the socket cleanly enough for the handshake to complete.

This is the **underdetermined (hypodox)** state in ``paradox_and_emergence_v003.md``'s vocabulary. Some granted principles apply (the socket constrains what icons would fit) but they leave the matter open (multiple candidates **could** fit, depending on what the brain happens to produce). The state has **slack**: the brain can guess, the wave-function-contribution of the prior facts can carry one guess to a fit, the same configuration could resolve in any of several ways depending on which firing pattern actually wins out.

The socket itself is best classified, on the refined anatomy of the scale, as a **delineated void** — a shaped, specified hole — sitting **within** an under-determined field of candidates (the delineated/undelineated split). The hole and its open surround are distinct: the socket is a definite empty shape (delineated), while the field of what-might-fill-it is the broader under-determined

region, and beneath both lies the undelineated void of what has not been shaped into a hole at all. The earlier framing that filed "socket open, candidate not yet emerged" simply under *under-determined* is sharpened by this split: the socket is the delineated void, the candidate-field is the under-determined surround.

This is the configuration in which most successful recall *can* succeed but has not yet. The brain is going to guess; one of the guesses will, if all goes well, fit; the configuration is open enough to permit success but not yet at success.

****Over-determined: too many candidates, partial fits, or over-specified socket.**** The brain is producing candidates that the mind is testing against the socket, but the configuration has gone wrong in one of three ways.

Too many candidates. Multiple icons are currently fitting the socket partially well; the system cannot settle on one because each has features the others lack. The mind is being pulled in several directions at once.

Partial fits with gaps. No candidate fits the socket fully; each candidate satisfies some of the socket's specifications and leaves others unfilled. The handshake cannot complete because no single icon is the icon being sought.

Over-specified socket. The mind has shaped the socket so narrowly that no icon the brain is producing can satisfy all the specifications. The mind is asking for too much; the brain is not finding any pattern that fits.

This is the **overdetermined (paradox)** state in ``paradox_and_emergence_v003.md``'s vocabulary. Granted principles apply incompatibly: the socket specifies, the candidates produce, the specifications and the candidates do not reconcile. The state forces overflow.

Two cognitive phenomena live in this state. **Tip-of-the-tongue** is over-determination of the partial-fit kind: the mind has sockets specified down to almost-the-right-name, with phonological and semantic constraints that almost match what the brain is producing, but no single candidate fits all the constraints. The candidates are partial-fitters; the gap is felt as the felt incompleteness of the tip-of-the-tongue state. **Trying-too-hard** is over-determination of the over-specified-socket kind: the mind has tightened the socket to require so much that nothing the brain can produce will satisfy it. The harder the mind tries, the narrower the socket; the narrower the socket, the less likely any candidate is to fit. The recall fails not because the brain lacks the parts but because the mind has refused to relax the shape.

The companion aphorism is from Diary 064B, 8 May 2026: **a*

balanced mind can remember easily*. Balance, here, is balance in the socket's specifications — neither so loose that anything fits (which is no recall but reverie) nor so tight that nothing fits (which is the over-specification just described). A balanced mind has its socket configured at the level where the brain's production has a reasonable chance of supplying a fitting icon. This is not an exhortation to relaxation in some general therapeutic sense; it is a structural observation about the mechanism. The four-state apparatus makes the observation precise: balance is the configuration in which the system is *under-determined* (open enough to admit a fit) rather than *over-determined* (closed against any fit).

****Fully-determined:** candidate fits socket, icon forms, recall succeeds.** The brain has produced a firing pattern whose icon fits the socket — the icon's points coincide with the socket's specifications, the icon is the icon the mind was asking for.

The handshake completes. The icon joins the collection of icons that the awareness-centre integrates; the centre re-centres on its new configuration; the cognitive task that initiated the socket is complete (or moves to its next stage, with a new socket shaped by the icon just acquired).

This is the *fully determinate* state in

`paradox\_and\_emergence\_v003.md`'s vocabulary. Granted principles apply and resolve the matter consistently and completely. The relationship has a definite character; the icon is the icon, the

recall has succeeded, the perception has formed, the deduction has reached its conclusion. The configuration is **just-settled** — not statically at rest but freshly arrived at the settled state, the live moment of resolution having just passed into the known (F10, F11). The over-determined press of the present click resolves **into** full determination as it settles; "fully-determined" names the just-settled state, not the live click itself. It stays settled until the awareness-centre, having shifted, calls for a new socket and the cycle begins again.

****The traversal, as a worked example.**** Perception is the clean case of the four states read as motion (D10). Take the icon of the pear forming. It begins **under-determined**: the activity layer is rearranging — firing patterns added and dropped, no whole yet closed, the field of candidates open (the **air** stage of the elemental loop, the rearrangement of parts). It **closes**: the configuration forces a single whole, the icon of the pear rather than a field of unbound firings — this is the making-of-one-whole, the point at which the socket is filled by exactly one candidate (the **water** turn). It is **lit**: the closed whole becomes the settled fact of the perception and is illuminated by the downward causation of §5, which is what gives it brightness (the **fire** stage). The icon does not **occupy** under-determination or full-determination as a fixed address; it **passes through** the scale from open rearrangement, to closure, to lit settled fact. This is what "motion along the scale" means at the level of mind, and it is the same passage the anchoring/creation split of §3.1

describes from the other side: the created dimensions are the under-determined pointers carried along the early part of the traversal, the shared dimensions the over-determined anchor the settled fact rests on.

4.4 Wave-function guess-and-confirm

The mechanism in §4.2 — **the brain doing its own thing**, with its firing patterns producing candidate icons that may or may not fit the current socket — is not as mechanical as a description of the brain's mechanics alone would make it. The brain is not just firing patterns randomly and waiting for one to fit. The fitting candidates are **made more probable** by the operation of the spirit-facts' wave-function-contribution.

Diary 059A, 17 January 2025, develops the two-stage **wave-function guess-and-confirm** account that explains this. Stage one: the brain guesses. It produces a candidate firing pattern on whatever basis its mechanics permit — context, ongoing activity, the residue of recent firings, partial cues from sensory input. Stage two: the prior memory's wave-function-contribution confirms or fails to confirm the guess. The facts that are stored, in the brain or in the spirit (the question of where memory lives is one the framework engages elsewhere), are part of the wave function that the brain's activity is collapsing under at this moment. Those facts make the firing patterns that would produce **the right icon** more probable than the firing patterns that would not.

This is what makes the brain's guessing not random. The guess is weighted by the facts that the spirit (or the long-term storage, or whatever the right name for the bearer of the memory turns out to be) makes available as wave-function-constraints. The brain does not have to find the right pattern by chance; the right pattern is preferentially produced because the wave function it collapses under is loaded by the prior facts.

This is the downward-causation mechanism of §5 below, viewed from the handshake side. The mind's facts (here, the spirit-facts that hold the memory) contribute to which layer-2 firing patterns are more probable, which means contributing to which layer-3 icons are more likely to emerge. The handshake completes because the brain's guessing is biased by the prior facts that the spirit makes available; the bias is wave-function contribution; the operation is downward causation.

The two stages are not temporally separable in any sharp sense. The brain's firing is continuous; the wave-function-collapsing is continuous; the guess and the confirmation are different aspects of the same ongoing process rather than two events in sequence. The diarist's "two-stage" language is a way of presenting the mechanism for understanding; the actual brain operates the two stages simultaneously, with the wave-function-contribution already loaded into the brain's activity at the moment the activity is occurring.

The framework's claim is that this is how recall **works**. The brain does not retrieve a memory from a stored representation in a file-cabinet sense; the brain produces a fresh firing pattern under wave-function-constraint, and the resulting icon at layer 3 fits the socket because the wave-function-constraint has loaded the brain's activity toward firing patterns whose icons do fit. Memory is not a file system; memory is the spirit-facts acting as wave-function-constraints on the brain's current guessing. The 064B **mind as socket, brain as parts-supplier** metaphor and the 059A **wave-function guess-and-confirm** mechanism are the same mechanism described from two angles.

4.5 The recentring of awareness

When the handshake completes — when a firing pattern produces a layer-3 icon that fits the current socket — the icon joins the collection of icons that the awareness-centre integrates. The collection is not the same after the new icon arrives. The centre re-centres.

This is where the centre-of-gravity machinery from ``gravity_and_the_past_v1_0.md` §13.6` applies in the icon case. The centre of a configuration of masses is **exactly where the configuration makes it be**; if the configuration changes (a mass is added, removed, or moved), the centre is in a new place. The centre is a virtual fact about the configuration; the configuration determines the centre completely. The same is

true at layer 3: the centre of an icon-collection is exactly where the configuration of icons makes it be. When a new icon joins the collection, the centre moves to its new configuration-determined location.

The recentring has consequences for what happens next. The next cognitive moment proceeds from the new centre. Whatever socket is created next is shaped by the new configuration of icons — the new centre may bias the socket toward dimensions the previous centre was not heavily biased toward, or may concentrate the socket in directions the previous centre was already concentrating in. The cycle continues: socket → guessing → wave-function-constraint → icon formation → recentring → new socket.

This is the dynamic of consciousness in the framework's sense: not a single stable state, but a continuous cycle of socket-creation and icon-formation, with the centre of awareness recentring at each completed handshake. The icons of awareness are not still; they are the icons currently in the centring, which is constantly being updated as new icons arrive and old ones fall out of the integration.

The 064A and 064B passages on the spirit's role in this cycle — *"memory is spirit is facts"* from Diary 060A, the **Spirit, DC*, and the *in-between state** passage of Diary 064B, 11 May 2026 — locate spirit's role at the wave-function-contribution step.

Spirit provides the prior facts that load the wave function; those facts make the brain's guessing biased toward icon-formation in directions the spirit's history supports. This is what makes the cycle not a random walk through icon-space but a structured walk, biased by the spirit-facts of the individual whose brain is doing the operating.

****The recentring is the Cartesian Theatre, reframed.**** There is no place in the brain where the icons are presented to an audience — no theatre, no homunculus in a seat. The framework grants Dennett that denial in full. But Dennett's denial of the **place** is not a denial of the **fact**. What there is, instead of a place, is an abstract configuration of mind: the centring that stages the most important icons — the most strongly and meaningfully connected — at its focus. Centring just **is** brightness by meaningful connection: the icon most connected to the rest is the brightest, is where the centre sits, is what gets "staged." The theatre is not a room; it is the configuration's focus, real on the same as-if basis a centre of gravity is real (F17). ``centre_of_awareness.md`` develops this with the centre-of-gravity-as-virtual-fact machinery of ``gravity_and_the_past_v1_0.md`` §13.6; here it is enough that the recentring of this section **is** the staging Dennett denied a place to, relocated from a place to a relational fact.

****Awareness maintains and replicates consciousness.**** The recentring does more than admit each new icon. The centre, by resting on its icons, **replicates** them — holds them recent, bright, and close,

refreshing the collection rather than letting it decay. This is why rehearsal aids recall: re-attending an icon replicates it, keeping it near the centre and bright, so that an imprecise later cue still finds a recently-attended icon waiting. Awareness is in this sense both the *wholeness* of mind — the integration that makes the icon-collection one — and its *heart* — the maintaining beat that keeps the collection live (F21).

4.6 The same mechanism in recall, perception, deduction, imagination

The unifying claim of this section is that recall, perception, deduction, and imagination are the same mechanism. They differ only in what shapes the socket.

****Recall.**** Socket shaped by the cue plus the spirit-facts whose wave-function-contribution is active. The brain guesses; the wave-function-contribution biases the guess; an icon fitting the socket forms; the recall succeeds. This is the case the mechanism was developed around.

****Perception.**** Socket shaped by the current sensory input plus the icons currently in the awareness-collection (which provide context). The brain guesses (in this case, the guessing is driven heavily by the immediate input from the senses); the wave-function-contribution from the spirit-facts also operates but is less load-bearing than in the recall case; an icon fitting the socket forms; the perception is constituted. The icon of the scene at hand is the icon that fits the socket the scene-

plus-context shapes.

That perception works by the same mechanism as recall, with the socket shaped differently, is consistent with the long-standing view in cognitive science that perception is **top-down** as much as bottom-up. The bottom-up component is the sensory input shaping the socket. The top-down component is the icon-collection context (and, via wave-function-contribution, the spirit-facts) shaping it further. The icon that forms is the icon that fits both components together. This is why we see what we expect to see, in the sense relevant to actual perception: the socket is shaped partly by expectation, and the icon that fits is the icon that satisfies both expectation and input.

****Deduction.**** Socket shaped by the icons already in play (the premises) plus the structural requirements of the inference. The brain guesses; the wave-function-contribution operates; an icon fitting the socket forms — and the icon, in this case, is the icon of the conclusion. Deduction is the same mechanism as recall, with the socket shaped by what-would-make-the-icons-already-in-play-add-up rather than by what-was-previously-true-of-the-individual. This is the framework's account of how deduction fits the more general account of cognition: it is not a separate faculty, but the icon-handshake-mechanism operating with a particular kind of socket-shaping.

The Log v6 **Significance, objectivity, and consensus** section

makes this point in the Log's vocabulary: *the same icon-machinery that produces ordinary awareness also produces deduction, simply by being aimed at an absence rather than a presence*. The substrate document's version sharpens what "aimed at an absence rather than a presence" amounts to structurally: the socket in the deduction case is shaped by what-would-make-the-premises-coherent rather than by what-is-being-recalled or what-is-being-perceived. The icon that fits this socket is the icon of the conclusion. Same handshake mechanism, differently configured socket.

****Imagination.**** Socket shaped less by external input (sensory or spirit-fact) and more by the imaginer's choices. The imaginer specifies the socket — *what if such-and-such were the case?* — and the brain produces an icon that fits the imagined specification. The wave-function-contribution from the spirit-facts still operates (which is why we cannot easily imagine things that contradict our deep beliefs), but the socket is freer than in the other cases.

That imagination operates by the same mechanism as recall is consistent with the corpus's repeated observation (Diary 047 *imagination-orientation-opposite-of-memory*; Diary 60B *imagination-as-holographic-mirror*) that memory and imagination are closely related. They are the same mechanism; they differ in the source of the socket-shaping. Memory's socket is shaped by the past (via spirit-facts); imagination's socket

is shaped by present specification (with the past constraining but not determining); perception's socket is shaped by present input plus context; deduction's socket is shaped by the icons in play plus the rules of inference.

These four cognitive functions are the same machinery in four configurations. This is what the unifying claim amounts to: recall, deduction, the handshake/socket/parts-supplier account, wave-function-guess-and-confirm, and the centre-of-awareness recentring are *the same mechanism viewed at different magnifications*. The mechanism is the handshake; the magnifications are the configurations of socket-shaping. This is what the substrate document delivers, and what the four working documents that follow can take as ground.

****The family is graded, and it includes sense.**** The functions above are not a list of four but stations on a graded family, indexed by how many icons stand on the *question* side of the handshake and how many on the *answer* side. At the simplest end is *sense* — the member the four-fold list omits and which belongs named here: a single incoming icon met by a single closure, minimal question, minimal search, the limiting case where the socket is shaped almost wholly by present input. Then *recall*: a cue-shaped question answered by one past icon. Then *deduction*: several premise-icons on the question side, the conclusion-icon as the answer. Then *imagination*: the question side specified by the imager, the answer side the most freely produced. What recall,

deduction and imagination share, and what sense has least of, is
will — the agency that shapes the socket deliberately; the family
runs from the will-less (sense) to the will-shaped (imagination),
the will being the common element across the upper members. The
family also sorts the *kind of fact* its answer is: a *learned*
fact (perception's answer, taken from the world), a *remembered*
fact (recall's, taken from the past), an *imagined* fact
(imagination's, specified rather than found). The imagined are not
failures of truth but *partially-untrue facts*: a fiction is a
partially-untrue fact — an icon with created components and a
worldly anchor weakened to nothing — which is exactly the fiction
limit-case the core-claim gloss named (F22).

4.7 Why this is the right way to apply the four-state apparatus

`paradox\_and\_emergence\_v003.md` §3 develops the four states of
determination as Virtualism's central machinery for analysing
configurations of granted-principles-and-fulfilment. The four
states are: non-determined (floor), under-determined (hypodox,
room for change), fully-determinate (steady state), over-
determined (paradox, force to a new level). The substrate
document's claim is that the handshake mechanism is precisely the
application of this apparatus to the mind-brain relation, with
the four states corresponding to the four configurations of the
socket-and-candidate relation.

This is the same apparatus that applies, in

`paradox\_and\_emergence\_v003.md`'s own worked cases, to the

emergence of new dimensions, the formation of compound entities, the resolution of paradoxes into new wholes. It is not a re-invention of the four-state apparatus for the cognitive case; it is the **same apparatus**, applied to a particular kind of relation (mind-layer socket meeting brain-layer parts-supply). The handshake mechanism is what Virtualism's four-state apparatus **looks like** when the relation in question is the mind-brain relation. The substrate document does not need to argue for the apparatus — that is done in ``paradox_and_emergence_v003.md``. The substrate document needs only to apply it, which is what §4.3 has done.

This is also what makes the unifying claim of §4.6 — that recall, perception, deduction, and imagination are the same mechanism — not a stipulation but a consequence. They are the same mechanism because they are the same four-state-apparatus operating on the same kind of relation (socket-and-candidate at layer 3), with the only difference being what shapes the socket. The four-state apparatus is dimension-blind; it operates wherever the configuration of granted-principles-and-fulfilment is in any of the four states. The cognitive cases are four configurations of its application; nothing more is required to unify them.

5. Downward causation as wave-function constraint

The handshake mechanism of §4 produces an icon at layer 3 when the brain's firing patterns are biased toward producing fitting

icons. The bias comes from the wave-function-contribution of the spirit-facts that the brain's local wave-function-collapses are constrained by. This section gives that mechanism in full, answers Kim's Causal Exclusion challenge, and completes the Supervenience sub-section that the 9–14 February 2026 essay-draft (Diary 063B) left as a header.

The framework's claim is that downward causation, in the sense required for psychological causation to be real, is wave-function constraint operating at the join between layer 3 (mind) and layer 2 (brain activity). The mind's facts contribute to which layer-2 firing patterns are more probable; the more-probable patterns are those whose layer-3 icons would fit the current socket; the icon that emerges is biased toward the icon the mind has been holding the shape of. This is downward causation in the strict sense: a fact at a higher level of organisation constrains what happens at a lower level, without requiring a separate causal pathway between the two.

Three clarifications fix what "downward causation" means here, so the term can be kept without the dualist baggage it usually carries. *First, it is the action of wholes on their parts* (F23). Downward causation is not a special mental force; it is what any whole does to its parts — the centre of gravity constrains the masses, the icon constrains the firings — and "downward" names a direction on the one constraint, not a second kind of causing. *Second, up and down are perspectives on a single wave-function constraint* (F6). The

brain constrains the mind (the parts produce the whole) and the mind constrains the brain (the whole biases the parts); these are not two causal transactions but one wave-function constraint seen from two ends. That two-endedness completes a **virtuous circle** (F9): the brain's activity produces the icon (the upward half), the icon biases the brain's next activity (the downward half), and the cycle of socket → guess → icon → recentring → new socket is this circle turning. Neither half is prior; the circle is the unit. **Third, the wholes that do this are emergent classical facts on a stability gradient, none of them absolutely permanent.** An icon is a stable whole, a rock a more stable one, a proton more stable still — yet even the proton may, on long enough timescales, decay; stability is a matter of degree, not an absolute (F23). This is the register to hold apart from Spirit's: the present, lit, boost-backed wholes are **stable-not-eternal**, whereas the eternal-fact register Spirit develops is what a whole becomes once it can no longer change — the two must not be run together, and §8 and `consciousness\_awareness\_spirit.md` keep the line.

5.1 Kim's challenge and its standard physicalist response

Jaegwon Kim's Causal Exclusion argument is the most pressing challenge to any theory of consciousness that wants to give the mind real causal work. Stated briefly: if the brain's activity is physically caused by prior physical states (the **causal closure** of the physical), and if mental states **supervene** on brain states (no mental difference without a brain difference), then either mental states are identical with brain states (and so the

mental does no causal work that the physical does not already do), or mental states are *excluded* from causation (because any candidate mental cause is preempted by the physical cause that underlies it).

Kim's is the modern, sharpened form of the oldest version of the same challenge: Princess Elisabeth's question to Descartes — how can an immaterial mind move a material body, when the two share no common currency in which cause could pass between them? Kim and Elisabeth are one challenge in two centuries' dress: Elisabeth asks how mind reaches matter at all, Kim asks how it could without being either redundant or excluded. The wave-function-constraint account answers both at once, because it supplies exactly the common currency Elisabeth found missing — facts, contributing to wave functions — without positing a second substance for the currency to cross between. §8 gives the Elisabeth answer in full; it is flagged here so that the reader sees the two challenges are met by one mechanism, not two.

The standard physicalist response is to embrace identity: mental states *just are* brain states; mental causation *is* brain causation under a different description. This avoids exclusion but at the cost of denying that mental states have their own causal contribution. Anti-reductionists have generally tried to preserve mental causation by qualifying causal closure, or by restricting supervenience to weak forms, or by denying that mental states would be excluded if they were genuinely distinct. None

of these moves has been entirely successful; Kim's argument continues to set the agenda of the debate.

The framework's response is structurally distinct from any of these. It accepts causal closure as commonly understood (the physical world is causally closed); it accepts supervenience as commonly understood (no mental difference without a brain difference); and it denies that the layer-3 facts are excluded, because the layer-3 facts contribute to the wave-function constraints that determine *which* of the physically possible brain states actually obtains. Causal closure does not require that the physically possible brain states are *uniquely* determined by physical antecedents; it requires only that the brain states are part of a closed physical chain. Quantum mechanics has long established that the closed physical chain is not uniquely deterministic — multiple outcomes are physically possible at each measurement event — and Virtualism's account of wave-function-collapse-as-Arc-selection (``change_and_present_v1_0.md` §4`) makes this precise: at each Arc, the universe selects one outcome from the wave function's structured possibility-space, with the selection biased by all the facts that are part of that wave function.

The layer-3 facts of the mind are among those facts. They are not excluded from the physical chain; they are part of what the physical chain selects from among. This is the structural shape of the framework's response, and it is precisely the shape that

the Diary 046 *Downward Causation Inset Day* passage developed in 2021 in the project's own terms: *DC functions by altering the quantum probabilities of photons moving (and therefore neurons firing)*.

5.2 The wave-function-contribution mechanism

The mechanism, stated in `change\_and\_present\_v1\_0.md`'s vocabulary as updated by the Iconism account: a quantum wave function is the application of all the relevant facts to the imminent future of a system. Wave-function-collapse at an Arc is the selection of one outcome from among those the wave function makes possible. The selection is biased by the relative probabilities the wave function assigns to each possible outcome. Those probabilities are themselves a function of *which facts are applied* in computing the wave function: more facts, more constraints, narrower distribution; fewer facts, fewer constraints, broader distribution.

The layer-3 facts of an individual's spirit are among the facts that apply in the vicinity of that individual's brain. They are not in some other dimension that the brain has no access to; they are facts about the configuration of layer-3 wholes that the brain's layer-2 activity is producing, and they are present as constraints on what wave functions can collapse in the brain's neighbourhood. The spirit-facts contribute by making the wave functions in the brain's vicinity *narrower* than they would otherwise be: the probability distribution over possible brain

states is biased toward the states whose layer-3 icons would satisfy the spirit-facts.

Most of the time, this contribution is small. The spirit-facts of an individual constrain the brain's wave-function-outcomes by making certain firing patterns more probable, but not by enough to override the local physical determinants of the brain's activity. The brain is doing its own thing, the wave functions collapse mostly according to the brain's local physical constraints, and the spirit-facts' contribution is too slight to divert otherwise-determined outcomes. This is most of cognition most of the time: the brain runs on its own dynamics, and the spirit-facts ride alongside without exerting a perceptible influence.

Occasionally, the contribution is enough to tip a balance. The wave function in the brain's vicinity is approximately balanced between two possible outcomes, both supported by the local physical determinants. The spirit-facts bias the wave function slightly toward one outcome rather than the other. The Arc selects that outcome. A specific firing pattern occurs that produces, at layer 3, an icon the mind was holding the socket for. The handshake completes. This is the moment-of-clarity, the butterfly's-wing case: the spirit-facts have done causal work, by biasing the wave-function-outcome toward the fitting icon.

The Log v6 *Spirit and fact-causation* section gives this in the

Log's own vocabulary: *the facts of an individual's spirit contribute to every wave function that collapses in the vicinity of their brain. Most of the time the contribution is too slight to divert outcomes that are otherwise heavily determined; occasionally — the moment of clarity, the butterfly's wing — the spirit is sufficient to tip a balance, and the outcome is swayed.* The substrate document keeps the Log's formulation and sharpens it by locating the mechanism at the layer-2-to-layer-3 join: the spirit-facts contribute by biasing which layer-2 firing patterns are more probable, which is to say by biasing which layer-3 icons are more likely to emerge.

5.3 The icon as virtual obstacle

The 22 October 2021 *Downward Causation Inset Day* passage of Diary 046 puts the mechanism in a vivid image: *icons as virtual obstacles affect photon chances*. The icon, at layer 3, is a virtual whole that constrains the configurations of layer-2 firing-patterns whose emergent wholes would be inconsistent with the icon's content. The constraint is real even though the icon is virtual, because the icon is a *fact* (a virtual fact, in Virtualism's sense — `foundations\_v001.md` §6) and facts contribute to wave functions (`change\_and\_present\_v1\_0.md` §4).

A photon traversing the brain — or, more generally, any quantum event that could go several ways — is governed by the wave function applicable in the brain's vicinity. The wave function

includes contributions from all the facts in play, which includes the layer-3 facts about which icons obtain. An icon that obtains is a virtual obstacle in the sense that the wave function must be consistent with that icon's continued obtaining; outcomes inconsistent with the icon are suppressed, outcomes consistent with the icon are favoured. This is the constraint operating.

The image of *obstacle* is apt because it makes the mechanism intuitive: a virtual obstacle does not block a photon's path the way a piece of glass does (a photon does not bounce off a virtual whole), but it shapes the photon's probability-distribution the way a quantum-mechanical potential does (the wave function incorporates the virtual whole's structural fact as a constraint on its support). The photon has more amplitude in regions consistent with the obstacle and less in regions inconsistent with it. This is what wave-function-constraint amounts to in the photon case, and the framework's claim is that the same mechanism operates wherever wave functions are computed in the brain's vicinity.

This is the answer to Kim's Causal Exclusion challenge. The layer-3 facts are not in competition with the layer-1 and layer-2 facts for a single causal role; they contribute *to the same wave function* that those lower-level facts also contribute to. The causal closure of the physical is preserved (the wave function is a physical object in the relevant sense; its constraints are physical constraints); supervenience is preserved (a difference

in layer-3 facts entails a difference in layer-2 firing patterns, because the layer-3 facts contribute to the wave function that determines the firing patterns); but the layer-3 facts are not excluded, because they are **part of** what the wave function incorporates, not an addition over and above what it already incorporates.

5.4 Completing the Supervenience sub-section

The 9–14 February 2026 essay-draft (Diary 063B) names Kim's three-part physicalist apparatus — causal closure, causal exclusion, supervenience — and develops responses to the first two. The Supervenience sub-section is a header without a body, flagged in the inventory for completion. This subsection completes the work.

The orthodox Kimian use of supervenience is to say: no mental difference without a physical difference. If two systems are physically identical (down to the last microphysical fact), they are mentally identical. The argument from supervenience to exclusion runs: if mental facts supervene on physical facts, then whatever causal work the mental facts seem to do is already done by the physical facts on which they supervene; the mental facts are therefore causally redundant.

The framework accepts supervenience as a constraint on the layered emergence: if two brains are physically identical at layers 1 and 2 (same connectomes, same firing patterns at every

instant), they have identical icons at layer 3. There is no floating layer-3 fact that fails to be grounded in some layer-2 fact. The layer-3 icon-and-mind layer is fully determined by the layer-2 firing-pattern layer; in this respect the framework is a supervenience theory.

What the framework denies is that supervenience implies exclusion. The denial rests on the wave-function-contribution mechanism. The layer-3 facts of the mind are not redundant with the layer-2 facts of the brain because they are not *additional* to those facts; they are facts *about the configuration of those facts*, and they enter the causal economy of the brain through the wave-function-contribution they make to the wave functions that the brain's quantum events collapse under.

This is not a wave function that excludes the brain's local physical determinants. It is the same wave function the brain would have anyway, with the layer-3 facts among its constraints. The brain's local physical determinants are themselves part of the wave function; the layer-3 facts are also part of it; the wave function is the joint object that incorporates both. When the wave function collapses, the outcome is one of the possibilities the wave function supports, and the bias toward that particular possibility is jointly determined by all the contributing facts. The layer-3 facts make a difference because they make the wave function different from what it would be without them; the difference shows up in the probability bias.

This is supervenience without exclusion. The mental supervenes on the physical (in the precise sense that fixing the physical fixes the mental); the mental is not excluded from causation (in the precise sense that the mental facts make a difference to what the physical does, by entering as constraints on the wave functions the physical collapses under). The orthodox Kimian argument runs from supervenience to exclusion via the assumption that supervenience implies redundancy; the framework rejects the assumption because the layer-3 facts are not redundant with the layer-2 facts on which they supervene, they are facts about *relations among the layer-2 facts* and so enter the wave function as constraints the layer-2 facts taken individually do not impose.

The diary's 14 February 2026 *Overdetermination* passage points at this without quite saying it: *the orthodox dialectic between reductionism and holism is unproductive because both share an unexamined assumption about Time — Time is taken as an extra, a given, in QM*. The unexamined assumption, on the framework's reading, is that the wave function is fully determined by the layer-1 and layer-2 facts independently of the layer-3 facts. With Virtualism's account of wave-function-as-application-of-relevant-facts ('change_and_present_v1_0.md' §4), the assumption falls. The wave function incorporates all the facts that obtain in the relevant region; the layer-3 facts obtain in the relevant region; the wave function incorporates them. Supervenience without

exclusion follows directly.

5.5 Bell's Inequality and the structural argument

The 9 February 2026 essay-draft's opening move is structurally load-bearing: *Bell's Inequality has falsified the EPR objection; this falsifies the deterministic reductionism that underwrites materialist denials of consciousness; any quantum experiment is already a worked example of downward causation, because the experiment itself causes the particle to take on the measured property.*

The framework's commitment to this move is structural. The materialist denial of consciousness typically presupposes a view of the physical world as a closed deterministic system in which higher-level facts have nothing to do. Bell's Inequality has empirically falsified the simplest forms of this view: the universe does not run on hidden variables in any local way. Whatever the right metaphysics of quantum mechanics turns out to be, it cannot be the one in which the physical world is exhaustively specified by local microphysical facts.

The framework's account of wave-function-collapse-as-Arc-selection (`change\_and\_present\_v1\_0.md`) is *one* metaphysical account consistent with the Bell-test results. It is not the only one — many-worlds, pilot-wave, and various other interpretations also accommodate the results. But all of them share the feature that local microphysical determinism is gone;

the wave function does *something*, and whatever it does is shaped by more than just the local microphysical state. The framework's claim is that the layer-3 facts of the mind are among the things that shape it.

This is the structural argument from Bell's Inequality to the non-exclusion of mental facts: the Bell-test rules out the metaphysical view that would underwrite Kim's exclusion argument; the alternative metaphysical views all permit higher-level facts to enter the wave-function-determination process; the framework's specific account of how they do so is the wave-function-constraint mechanism of this section. The argument is not that Bell's Inequality entails the framework. It is that Bell's Inequality rules out the strong reductive metaphysics that the exclusion argument requires, leaving the framework's account as one available alternative.

Any quantum experiment, on the framework's reading, is a worked example of downward causation: the experimental apparatus (macroscopic, layer-2-or-higher in its own internal structure) constrains the particle's wave function in ways that bring about specific outcomes that depend on the apparatus's configuration. The apparatus is doing downward causation; the framework's claim is that the brain, which is similarly a macroscopic structure configured in highly specific ways, does the same kind of work on the wave functions in its neighbourhood, with the layer-3 facts of mind playing the role the apparatus's configuration

plays in the experiment.

5.6 What this gives the framework

The wave-function-constraint account of downward causation gives the framework four things.

****It gives a structural answer to Kim.**** Causal closure preserved, supervenience preserved, exclusion denied — through the specific mechanism by which layer-3 facts enter the wave-function determination of layer-2 outcomes. The answer is not a denial of Kim's premises but a structural account of how those premises can both obtain without their conjunction yielding exclusion.

****It gives a positive mechanism for psychological causation.****

The mind's facts do real causal work on the brain. The work is not separate-substance interaction (no Cartesian interaction problem); it is wave-function-contribution. The mind influences the brain because the mind's facts are among the facts the brain's wave functions answer to. The Log v6 **Spirit and fact-causation** section's formulation — **the mind influences the brain because the facts of the mind are among the facts the brain's quantum processes must answer to** — is precisely this.

****It gives a unified account of the handshake mechanism of §4.****

The wave-function-guess-and-confirm of §4.4 is the same mechanism as the wave-function-constraint of this section; the former describes the mechanism from the cognitive side (brain

guesses, prior facts confirm), the latter describes it from the metaphysical side (mind's facts bias the wave functions in the brain's vicinity, biased wave functions produce fitting firing patterns more often). The two are not separate accounts but the same account viewed at different levels.

****It gives the framework a positive role for the spirit.**** Spirit is the collection of all true facts about an individual (Log v6 **Consciousness, awareness, spirit**). The spirit's facts contribute to the wave functions in the brain's vicinity. This is how the spirit influences the conscious self: not through some separate mental substance, but through the contribution the spirit's facts make to the wave-function-determination of brain outcomes. The Iconism-Spirit boundary, in this respect, is precisely defined: the **mechanism** of psychological causation (wave-function constraint) is in Iconism's territory; the **ontology** of the spirit-facts (eternal, persistent, accumulated over a lifetime, perhaps surviving demergence) is in Spirit's territory. The two territories meet at the wave-function-contribution step.

5.7 What the framework does not claim about wave-function constraint

It does not claim that this mechanism has been empirically demonstrated in any specific brain at any specific moment. The claim is structural: this is how downward causation **would** operate consistently with quantum mechanics and with Virtualism's account of wave-function-as-application-of-facts. Whether any

specific clinical or experimental phenomenon is best explained by this mechanism is for empirical work to determine.

It does not claim that wave-function-constraint is the **only** form of downward causation. The 21 May 2026 addendum passage on the **four flavours of downward causation** via the Mandala structure (Fire→Earth→Air→Water→Fire) suggests there are multiple kinds of DC operating in the framework. The wave-function-constraint mechanism developed here is the kind relevant to psychological causation; other kinds may operate in other domains (e.g., the centre-of-gravity's constraint on its parts is a kind of DC, but it operates classically rather than through wave-function-collapse). The framework permits multiple kinds of DC; this section develops only the kind relevant to the substrate document's spine.

It does not claim that the wave-function-constraint mechanism explains **every** feature of psychological causation. Some features — the qualitative phenomenology of agency, the felt character of choosing one option over another, the experience of deliberation — are not directly explained by the mechanism even though they are consistent with it. The framework offers the mechanism as the structural account of **how** psychological causation can be real; the phenomenology of the cases is material for ``centre_of_awareness.md`` and ``consciousness_awareness_spirit.md`` to develop further.

6. Worked applications

The handshake mechanism of §4 and the wave-function-constraint account of §5 are stated abstractly. This section shows the mechanism operating in particular cases. The cases are chosen because they illustrate the framework's claims under the kind of phenomenological detail that any working theory of cognition has to be able to handle. They are: confirmation bias, tip-of-the-tongue, and the four cognitive modes (recall, perception, deduction, imagination) treated through their socket-shaping differences.

The point of the worked applications is not to **prove** the framework — none of these cases is a knockdown argument. The point is to **exhibit** the framework: to show what the substrate mechanism looks like when run on familiar cognitive phenomena, so that the working documents that follow can refer back to these illustrations rather than re-developing them.

6.1 Confirmation bias

Confirmation bias is the tendency to register evidence consistent with what one already believes more readily than evidence that challenges those beliefs. In the framework's vocabulary: the icons currently in the awareness-collection bias which firing patterns the brain produces in response to new input, and the bias favours patterns whose layer-3 icons are **consistent with**

the existing collection.

The mechanism is the wave-function-guess-and-confirm of §4.4 operating in a configuration where the prior icon-collection is heavily weighted in the dimensions where the new input is being assessed. The socket created in response to the new input has its shape biased by the existing icons; the wave function that the brain's local activity collapses under is biased by the existing icons' facts (which are layer-3 facts and so contribute to the wave function, per §5); the firing pattern that is more probable is the one whose icon would fit the socket-as-biased; the icon that forms is the icon consistent with the existing collection.

Confirmation bias is therefore not a flaw in cognition but a direct consequence of the mechanism by which cognition works. Recall and perception are both biased by the prior icon-collection because the prior icon-collection is part of what shapes the socket and part of what contributes to the wave function the brain's activity is constrained by. The bias is the working of the mechanism; the same mechanism that produces successful recall produces confirmation bias when the input is partially inconsistent with the prior collection. The two phenomena are not two mechanisms but one mechanism in different configurations.

This is part of why confirmation bias is *hard to overcome*. It is not a layer of cognitive work added on top of perception or

recall that one could decide to do without; it is the way perception and recall actually work. To overcome confirmation bias requires *deliberately reshaping the socket* to be less biased by the existing icon-collection — which can be done, but which requires the kind of effort that any reshaping of the socket requires. The mechanism of overcoming confirmation bias is the same mechanism as the mechanism of confirmation bias; the difference is in what the mind is making the socket open to.

The framework's claim is consistent with what cognitive psychology has found empirically about confirmation bias: it operates pre-consciously, it is hard to introspect, it is context-dependent in ways that match the framework's prediction about which icons are most heavily in the awareness-collection at a given moment. The framework does not predict the specific empirical details — those are for cognitive psychology — but it predicts the *kind* of phenomenon confirmation bias would be: not a discrete error-process but a feature of the basic mechanism.

6.2 Tip-of-the-tongue

The tip-of-the-tongue phenomenon — the felt presence of a word or name that one cannot quite produce — is over-determination of the partial-fit kind, in the vocabulary of §4.3. The mind has specified a socket with rich constraints: the name has roughly this length, starts with roughly this letter, refers to roughly this person, has roughly this resonance with other names one

knows. The brain is producing firing patterns whose icons partially satisfy these constraints but not fully — the icon *fits some of the socket* but leaves *some of the socket unsatisfied*. The handshake cannot complete; multiple candidate patterns are partial fitters; the configuration is over-determined in the sense that the candidates do not jointly satisfy the socket's specifications and no single candidate satisfies all of them.

The felt phenomenology of tip-of-the-tongue — the sense that the word is *there*, almost-but-not-quite — is the phenomenology of the over-determined configuration. The mind has the shape and some of the content; what is missing is the specific point at which a single icon would complete the handshake. The mind is aware of the shape (which is why the experience is felt as *the word is there*) and aware of the missing piece (which is why the experience is felt as *not quite there*). The felt-ness of the absence is precisely the felt-ness of the unfilled parts of the socket.

The framework predicts two ways the tip-of-the-tongue configuration resolves. *First, by waiting.* If the mind relaxes its specifications — opens the socket somewhat by relaxing the most peripheral constraints — the configuration moves from over-determined to under-determined. The brain's continuing activity can then produce a firing pattern whose icon fits the relaxed socket. This is the experience of the word coming back

when one has stopped trying: the relaxation of the socket was what permitted the handshake to complete.

Second, by changing the cue. If the mind shifts the socket in some non-trivial way — recalls additional context, brings a different cue to bear — the socket changes shape, and the brain's continuing activity may now produce a candidate that fits the new shape. This is the experience of suddenly recalling the word when one has thought of something else: the shift in cue changed the socket; the candidate that the brain was already producing now fits.

The 064B aphorism *a balanced mind can remember easily* is the companion observation. A mind that is too tense — specifying the socket too narrowly, leaving no room for the brain's candidate to fit — is in over-determination; the recall is blocked by the over-specification. A mind that is too loose — specifying the socket too broadly — is in under-determination but with such broad shape that almost any candidate fits and none is specifically *the* recall. The balance is the configuration in which the socket is open enough to admit a fit and specified enough that only the *right* candidate fits. The framework's claim is that this is a structural fact about the mechanism, not a piece of folk-psychological advice.

6.3 Recall as wave-function guess-and-confirm

The Diary 059A treatment of recall, 17 January 2025, is the most

careful single account in the corpus. It is worth working through in the substrate document's vocabulary, because it shows how the wave-function-constraint mechanism produces the phenomenology of successful recall.

****One correction frames the whole section.**** An earlier reading let wave-function language run across the **whole** of recall, which invited the fair objection that the framework had a "shouting wave function" doing implausibly much work at a distance. The fix is to see that recall is **two** steps, and that wave-function language belongs to only one of them (F13).

The first step is a ****forceless Leibniz connection****. The socket is an answer-shaped void; the matching past-icon fills it **by Leibniz's Law**, because the void and the icon are indiscernible at the relevant points and so are, at those points, the same. This filling is forceless — there is no carrier, no signal, no energy crossing a gap, and crucially no dependence on distance: a past-icon however remote in time fills its matching void as directly as a recent one, because sameness-by-Leibniz is not a force that attenuates but an identity that simply holds (the forceless icon-connection). Nothing here is a wave function; nothing here is guessed.

The second step is the ****brain's guess-and-paint****, and **this** is where wave-function language properly stays. The brain, unbalanced and seeking, produces firing patterns — it guesses — and brightens the ones that begin to fit, painting the icon in by degrees until a

pattern's icon registers the fact the void was shaped for. The wave function is the structure of **this** guessing in the brain's vicinity, biased by the spirit-facts; it is local, it is the brain's, and it does not have to shout across time because the Leibniz connection has already done the connecting. Reserving wave-function talk for the brain step is the whole of the fix.

Two riders follow. The brain may ***settle on an apparent fit*** (F7): a pattern whose icon fits the socket **well enough** can complete the handshake even when it is not the icon sought — which is the mechanism of misremembering and of confident error, the socket filled by a near-match the brain brightened too soon. And the completing handshake ***is a local Arc*** (F12): the moment the void fills and the icon lights is a selection-event of exactly the kind ``change_and_present_v1_0.md` §4` describes, a present clicking into the settled past — recall is not a read from a store but an Arc resolving in the brain's neighbourhood. The relaxed-self corollary of §4.3 reappears here with a mechanism: a tense or over-specified self mis-shapes the void, so that the right past-icon no longer matches it and the forceless connection cannot fire; **a balanced mind remembers easily** because it lets the void keep the shape the sought icon actually has.

A name is sought — a friend's name, say, in a context where it would normally be available but is not currently coming. The cue is the context: this person, this situation, this set of properties one knows about them. The socket is shaped by the cue

and by whatever spirit-facts about this person are accessible to the wave-function-contribution mechanism.

The brain is doing its own thing. Firing patterns are forming.

Most of them are not patterns whose layer-3 icons would fit the socket. The wave-function-contribution from the spirit-facts is operating in the brain's vicinity, biasing the probabilities of each pattern. Patterns whose icons would produce the friend's name are more probable than they would otherwise be; the bias is small in any single moment but accumulating across the brain's continuing activity.

At some point, a firing pattern emerges whose icon fits the socket. The icon contains the name. The handshake completes. The friend's name is recalled.

The framework's claim is that this is **all** there is to recall.

There is no file-system in the brain that the recall is reading from. There is no memory-store with the name held in a fixed location. There is a brain doing its activity, biased by the wave-function-contribution of the relevant spirit-facts, and at some moment the brain produces a firing pattern whose icon contains the name. The "stored memory", in the framework's account, is the spirit-fact whose wave-function-contribution biased the brain toward producing the name-containing pattern.

The stored memory is not in a brain location; it is in the spirit, and it acts on the brain through wave-function-

contribution.

This is what Diary 060A, 8 April 2025, captures in **memory is spirit is facts**. Memory just **is** the spirit-facts contributing to the brain's wave-function-collapses. There is no separate memory-system in the brain that holds the facts; there is the brain doing its activity under wave-function-constraint, and the spirit-facts are what supply the constraint. The neural correlates of memory — the synaptic strengthenings, the hippocampal involvement, the consolidation processes — are the brain's adaptations **to** the wave-function-constraint, not the substrate of the constraint itself.

If the spirit-facts are all present as constraints, a question presses: why is recall **effortful** — why is the rest of the spirit not simply available at once? The answer is ***siloing, and siloing is a brightness effect*** (F19). The living brain is loud: its own activity is so bright that it walls the currently-lit bundle off from the vast quiet field of spirit-facts, admitting only what the guess-and-paint step manages to brighten into the socket. The brain's brightness is not only what lights an icon; it is what **silos** the mind, keeping the lit bundle separate from everything not currently being brightened. This is why the field is not open all at once, and why a quieter brain (the relaxed self, the hypnagogic state) lets more through. The consequence past the brain's brightness ceasing altogether — what happens to the siloing when the loud brain is gone — is a Spirit-project matter and is

flagged, not developed here: it is the post-mortem un-siloing question, and it belongs with the eternal-fact ontology, not with the in-life mechanism.

This is a substantive theoretical claim about what memory is, and it has consequences for how the framework handles neurodegeneration. The 24 February 2026 passage of Diary 063B develops one: *Alzheimer's as loss of retrieval, not loss of fact*. The damaged brain loses the ability to produce firing patterns whose icons fit the socket the recall would require; the spirit-facts are unchanged (memory in the framework's sense is preserved), but the brain's ability to register them through the wave-function-contribution-and-pattern-production mechanism is impaired. This is consistent with the clinical observation that Alzheimer's patients sometimes recall in unusual conditions or under unusual cuing — the recall is not lost in the sense of "the fact has gone"; it is lost in the sense of "the brain is no longer producing the patterns that would register the fact". The framework's account of memory predicts this kind of dissociation; a file-system account of memory does not.

6.4 Perception under the same mechanism

Perception is the handshake mechanism with the socket shaped by current sensory input plus the icons currently in the awareness-collection. The brain is doing its own thing; the wave-function-contribution from the spirit-facts is operating; an icon fitting the socket forms.

Two features distinguish perception from recall in the framework's account. First, the socket-shaping is dominated by sensory input rather than by stored facts. The contribution from spirit-facts is still operating (which is why perception is context-dependent and influenced by expectation), but the load-bearing component of the socket is the sensory input. Second, the brain's guessing is more constrained by the sensory input — the firing patterns the brain produces in perception are heavily driven by the input from the senses, with the higher-level cognitive contributions operating as bias rather than as primary source.

The framework's account is consistent with the cognitive-science finding that perception is top-down as well as bottom-up. The bottom-up component is the sensory input shaping the socket. The top-down component is the icon-collection-context (and the spirit-facts via wave-function-contribution) shaping the socket further and biasing which firing patterns are more probable. The icon that forms is the one that fits both components together — which is why we see what we expect to see, in the precise sense that perception's socket is shaped by expectation, and the icon fitting the socket is therefore the icon expectation has shaped the socket toward.

The framework also gives an account of perceptual illusions and hallucinations. A perceptual illusion is a configuration in

which the sensory input and the prior icon-collection together shape a socket that admits an icon not matching the actual worldly object. The icon that forms is the icon the socket admits; the perception is non-veridical because the socket was shaped by mistaken context, not because the mechanism malfunctioned. A hallucination is a configuration in which the socket is shaped almost entirely by prior icon-collection (and internal state) with very little sensory input contribution; the brain produces an icon that fits the internally-shaped socket without external grounding. Both phenomena are the same mechanism as veridical perception, configured differently.

This is the framework's account of *seeing what is not there* and *not seeing what is there*: not a separate mechanism, but the same socket-and-handshake mechanism operating with the socket shaped in ways that depart from veridical configuration. The mechanism is the same; the configuration is what produces the phenomenology.

6.5 Deduction as icons-fitting-an-abstract-hole

Deduction, in the framework's account, is the handshake mechanism with the socket shaped not by sensory input or by a recall cue but by the icons already in play (the premises) plus the structural requirements of the inference. The Log v6 *Significance, objectivity, and consensus* section captures this: *the same icon-machinery that produces ordinary awareness also produces deduction, simply by being aimed at an absence rather

than a presence*.

The framework's sharper version: deduction's socket is shaped by *what would make the icons in play coherent*. The premises are icons currently in the awareness-collection. The conclusion-to-be-derived is the icon that would complete the collection in a way the structural requirements of the inference permit. The mind specifies the socket as the shape of the icon that would make the premises lead to the relevant conclusion; the brain produces firing patterns under wave-function-constraint; an icon fitting the socket emerges; the deduction succeeds.

This account has the virtue that it explains why deduction can fail in the same ways recall can fail. *Over-specified socket*: the mind has demanded too tight a constraint on the conclusion and the brain cannot produce a firing pattern whose icon satisfies. *Under-determined socket*: the premises do not constrain the conclusion tightly enough and multiple icons could fit; the deduction is indeterminate. *Non-determined*: no deduction is being attempted; the brain produces ordinary icons without the deduction-socket-shaping in operation. The same four states of §4.3 apply; the same mechanism is operating.

The account also explains why deduction is *easier* when the relevant icons are already strongly in the awareness-collection. The wave-function-contribution from those icons biases the brain toward producing further icons that cohere with them; the

deduction-socket is shaped by the icons in play and is therefore biased by them in the wave-function-contribution sense. A familiar inference proceeds easily because the wave-function-contribution from the prior icons is strong; an unfamiliar inference proceeds with difficulty because the wave-function-contribution is weaker and the brain has to produce candidate firing patterns without much guidance. This is consistent with the cognitive-science finding that expertise — having many relevant icons strongly in the awareness-collection — makes domain-specific deduction much easier than the same inferences would be for a novice.

The framework also makes sense of why deduction can take time. The handshake does not complete instantaneously; the brain has to produce a firing pattern whose icon fits the socket, and the brain's production rate is finite. The wave-function-contribution biases the process but does not bypass it. Difficult deductions take time because the socket is tightly specified, the candidates are partial fitters, and the brain has to keep producing until a fitting pattern emerges. This is the **thinking-hard** phenomenology: the mind holds the socket, the brain produces candidates, none fits, the brain produces more, until one fits.

The 14 February 2026 **Overdetermination** passage of Diary 063B hints at this in its closing observation that **facts of the past* (recalled memories) compete with realities of the present, and

where a fact is recalled and understood it can become a new constraint on brain behaviour*. The recalled facts contribute to the wave-function-constraint; the recall thereby creates a new constraint; deduction works partly because recall has put the relevant facts in play as wave-function-contributors. The mechanism is unified across the cognitive functions because it is the same mechanism.

6.6 Imagination

Imagination is the handshake mechanism with the socket shaped by the imaginer's own specifications rather than by external input, recall cue, or deductive structure. The imaginer says: *what if such-and-such were the case?* The mind shapes a socket corresponding to *such-and-such*; the brain produces firing patterns under wave-function-constraint; an icon fitting the socket forms; the imagination is constituted.

The framework's account predicts two features of imagination.

First, imagination is biased by what we already believe. The wave-function-contribution from the spirit-facts operates in imagination as in any other handshake; we cannot easily imagine things that contradict our deep beliefs because the wave-function-contribution from those beliefs biases against the firing patterns that would produce icons of the contradiction.

This is consistent with the well-attested difficulty of imagining genuinely alien cultural arrangements, genuinely alien physical possibilities, and genuinely alien moral configurations:

the wave-function-contribution from one's existing icon-collection makes the alien icon hard for the brain to produce.

Second, imagination has a degree of freedom that the other cognitive modes do not. Because the socket is shaped by the imaginer's specifications and not (or not heavily) by external input, the imaginer can specify the socket in directions the other cognitive modes do not have access to. The imaginer can choose to imagine *this* configuration rather than *that*; the mind can specify what the socket is open toward. The brain has to produce a firing pattern that fits, and the wave-function-constraint still operates, but the socket-shaping is more freely available than in perception or recall.

The framework's account is consistent with the corpus's repeated observation that memory and imagination are closely related. Diary 047 (16 February 2022) — *imagination has an orientation opposite to memory*; Diary 60B (May 2025) — *imagination-as-holographic-mirror*; Diary 064B (12 May 2026) — *the mind creates a socket for the brain to fill or fuck over*, which is the mechanism for both. They are the same mechanism; they differ in which way the socket-shaping flows. Memory's socket is shaped by the past acting through wave-function-contribution; imagination's socket is shaped by the imaginer specifying forward, with the past still operating as wave-function-contribution constraint but not as the source of the socket-shape.

6.7 Summary of the worked applications

The same handshake mechanism — socket created at layer 3, brain producing firing patterns at layer 2 under wave-function-constraint, icon emerging at layer 3 when a firing pattern's icon fits the socket, awareness recentring on the new icon — is operating in all of:

- Confirmation bias (socket biased by existing icon-collection, candidates biased toward consistent icons).
- Tip-of-the-tongue (socket over-specified, candidates partial fitters).
- Recall (socket shaped by cue and spirit-facts, wave-function-contribution from spirit-facts biases the candidates).
- Perception (socket shaped by sensory input and context, wave-function-contribution operates as expectation-bias).
- Deduction (socket shaped by icons-in-play and inference rules, wave-function-contribution biases toward coherent further icons).
- Imagination (socket shaped by imaginer's specifications, wave-function-contribution operates as belief-bias).

In each case the **mechanism** is the same; the **configuration** of socket-shaping is what differs. This is the unifying claim of the substrate document — recall, deduction, the handshake/socket/parts-supplier account, wave-function-guess-and-confirm, and downward causation via wave-function constraint are the same

mechanism viewed at different magnifications — made operational through these worked cases.

The unification is not merely descriptive. It is what allows ``icon_and_world.md``, ``consciousness_awareness_spirit.md``, ``centre_of_awareness.md``, ``conceivability.md``, and ``ai_and_sentience.md`` to refer to a single underlying mechanism when their developments require it. Without the unification, each of the working documents would have to develop its own account of the cognitive functions it engages; with the unification, all five share the substrate worked out here. The economy of the substrate document is what its substrate role requires.

7. Continuity of consciousness across Arcs

Diary 059A (12 February 2025) poses what the framework recognises as the **second hard problem**: not only the standard qualia-from-physical-activity question, but also a **self-from-qualia** question — how does a single conscious life persist through the continuous selection-events the framework calls Arcs? The universe selects continuously (``change_and_present_v1_0.md` §2`); an observer's brain is part of the universe doing the selecting; what is the structural account of how a single consciousness obtains continuously across those selection events?

The corpus does not contain a finished answer to this question.

Diary 064B (3 May 2026) offers a partial articulation using Claude's restartability as an analogue for sleep-and-waking; the positive treatment is still to be written. This section offers the positive sketch the substrate document is positioned to give, in the layered-emergence vocabulary developed above. The sketch is not claimed as a final resolution; it is what the framework can say at this stage of the substrate work. The working document `centre\_of\_awareness.md` may refine the sketch with its fuller treatment of the awareness-centre. Spirit-side material (what happens to consciousness in death, in the in-between state, in re-emergence) is for the eventual Spirit project to develop.

7.1 The question precisely posed

The Arc, on Virtualism's account (``change_and_present_v1_0.md` §2`), is the selection-event by which one outcome from a wave function's structured possibility-space is converted into the present. The universe is producing Arcs continuously, at every quantum-mechanically significant event throughout its volume. An observer's brain is in this volume; the brain is producing Arcs as a natural consequence of being a physical system in the universe. The brain's firing patterns are themselves the result of many Arcs being selected within the brain at any given moment.

The question is what makes the conscious life of a single observer **continuous** across these many selections. At layer 1

(the material brain), the answer is straightforward in neuroscience's vocabulary: the material brain has continuity in its chemistry and its tissue; neurons fire and reset on millisecond timescales while remaining the same neurons across the firing-reset cycles. At layer 2 (firing patterns), the answer is more subtle but still tractable: a particular firing pattern is short-lived, but the connectome and the ongoing chemical context support the production of **related** firing patterns at successive moments, giving continuity of activity even though no specific pattern persists.

At layer 3 (icons and mind), the question is sharper. Icons can form and dissolve; the awareness-centre can move; the configuration at one Arc is not the configuration at the next. What makes the **same** consciousness persist across the continuous formation, dissolution, and re-formation of icons? The framework needs an answer that does not collapse into either (a) reducing layer-3 continuity to layer-1 or layer-2 continuity (which would deny that layer 3 has its own dynamics), or (b) positing a separate substantial soul that persists independently of layers 1 and 2 (which would re-introduce the Cartesian interaction problem).

The framework's answer, in the layered-emergence vocabulary, locates the continuity precisely at layer 3.

7.2 The positive sketch

Continuity of consciousness across Arcs is continuity of the third-layer wholes (icons and the awareness-centre) across the continuous selection events, even as the second-layer firing patterns and the first-layer material brain are changing under the same selection process.

The mechanism that produces this continuity has three components.

****First, the wave-function-contribution mechanism of §5 operates continuously.**** The spirit-facts of an individual contribute to the wave functions in the vicinity of that individual's brain at every moment. The contribution biases the brain's Arc-selections toward firing-patterns whose icons are **consistent with** the existing spirit-facts. As new icons form, they accumulate as spirit-facts; the spirit grows continuously; the wave-function-contribution at each moment includes the spirit-facts accumulated through every preceding moment. The continuous accumulation of spirit-facts is what biases the brain's continuous Arc-selections toward producing firing patterns continuous with the prior icon-configuration. Continuity at layer 3 is therefore not despite the Arc-selection process; it is **produced by** the Arc-selection process under the wave-function-contribution of the prior facts.

****Second, the centre of awareness recentres at each completed handshake but does not jump discontinuously.**** The centre of a configuration of icons is, by §3 and by
`gravity\_and\_the\_past\_v1\_0.md` §13.6, exactly where the

configuration makes it be. When a new icon arrives, the configuration changes, and the centre moves to its new configuration-determined location. But the new icon is biased toward consistency with the existing collection (by the wave-function-contribution); the centre's new location is biased toward proximity to its old location (by the new icon being biased toward consistency). The centre's trajectory is continuous because the icon-configuration is biased continuous. This is the structural account of the felt continuity of the self-as-aware: not a separate self-substance persisting through change, but a centre whose location at each Arc is constrained to be close to its location at the previous Arc by the wave-function-contribution mechanism.

****Third, the icons themselves can persist across many Arcs.**** An icon is a layer-3 fact that obtains when a firing pattern at layer 2 produces a whole at layer 3 indiscernible from the worldly whole the icon is iconic of. While the firing pattern that supports the icon at one Arc may not be the same firing pattern that supports it at the next Arc (because the brain's firing patterns are short-lived), the **kind** of firing pattern that supports the icon may be reliably reproduced across many Arcs (because the connectome and the wave-function-constraint mechanism bias the brain toward producing icon-supporting patterns of the same kind). The icon-as-virtual-whole therefore persists across many Arcs, even though the layer-2 firing-pattern that supports the icon at any one Arc does not. This is the

multiple-realisation point of `icon\_and\_world.md`'s topic, extended across time: the icon is individuated by its layer-3 relation (indiscernibility-with-worldly-whole), and the layer-2 substrate that supports it can vary from Arc to Arc without changing the icon. Continuity of the icon across Arcs is continuity of the layer-3 relation, not of the layer-2 substrate.

These three components — continuous spirit-fact accumulation biasing continuous Arc-selection, centre-of-awareness recentring without discontinuous jumping, icons persisting at layer 3 across changing layer-2 substrates — jointly account for the continuity of consciousness across the continuous Arc-selection process.

7.3 Why this is not a discontinuity but a continuous flow

The framework's response to the second hard problem is, in structural shape, the same kind of response Virtualism gives to the standard hard problem. Consciousness is not a property added on top of physical activity; consciousness is what physical activity *amounts to* at the layer where indiscernibility with worldly wholes can hold. Similarly, continuity of consciousness is not a property added on top of the continuous physical process; continuity of consciousness is what the continuous physical process *amounts to* at the layer where the wave-function-contribution-of-prior-facts biases the Arc-selections to be continuous with the prior icon-configuration.

There is no extra fact required to explain the continuity. The Arc-selection process is continuous (Arcs occur continuously, one after another, with no gaps in time); the wave-function-contribution of prior facts is continuous (the prior facts accumulate without loss, contributing to every wave function in the relevant region); the icon-configuration at each Arc is biased toward consistency with the prior configuration. The result is that the icon-configuration evolves continuously, the awareness-centre moves continuously, and the spirit-facts accumulate continuously. Continuity all the way down — or rather, continuity all the way up, since the framework is layered-emergent and the higher layers inherit continuity from the constraints the lower layers and the wave-function-contribution together impose.

This is consistent with the felt continuity of consciousness in ordinary experience: we feel ourselves to be the same person moment to moment, the icons we are aware of carry over from one instant to the next, the centre of awareness has the same character even as it moves. The framework's claim is that this felt continuity is not an illusion produced by some failure of introspection; it is the accurate registering of the structural continuity that the wave-function-contribution mechanism produces.

7.4 Sleep, dreams, and discontinuity

The framework also has to give an account of the discontinuities

in consciousness. Sleep is the standing example: the felt character of consciousness is interrupted between going to sleep and waking; dreams are felt as not-quite-the-same-kind-of-consciousness as waking; some sleep is felt as no consciousness at all. The framework's account of continuity has to be compatible with these discontinuities.

The Diary 058 (27 November 2024) *Consciousness* passage develops the picture in the corpus's own vocabulary: *sleep as centre moving away from icons; dreams as icons bidden by spirit*. In the substrate document's vocabulary: sleep is a configuration in which the wave-function-contribution from current sensory input is much reduced (the senses are largely shut down, the brain is in a different metabolic state) while the wave-function-contribution from spirit-facts continues. The brain continues to produce firing patterns under wave-function-constraint, but the constraint is dominated by spirit-facts rather than by the sensory-input-shaped socket that ordinary perception provides.

Dreams, on this account, are the icons that form when the spirit-facts dominate the socket-shaping in the absence of sensory input. The icons are *icons bidden by spirit*: they fit sockets shaped by accumulated facts rather than by current input. The phenomenology of dreams — vividness, emotional salience, occasional incoherence at the macro-level despite local coherence — matches what the framework would predict: icons formed under spirit-fact-dominated wave-function-

contribution, without the corrective discipline of sensory input to shape the socket toward veridical-perception.

The corpus's **icons bidden by spirit** names one kind of dream — the ***mind-origin*** kind — not the whole of dreaming; the sharper account distinguishes dreams by ***origin, with agency as the clue*** (D6). The through-line first: **all** dreams are experiences of ***awareness***. Awareness does not switch off in sleep; it is continuous through the night, and what varies between waking and dreaming is not whether awareness obtains but the **origin** of the icons it centres on and the **agency** it has over them (F16, F21). Two origins divide the dreams. ***Brain-echo*** dreams are low-agency: the brain, decoupled from sensory discipline, throws up firing patterns whose icons awareness merely witnesses — the macro-incoherence of ordinary dreaming being the signature of a brain guessing without correction. ***Mind-origin*** dreams are higher-agency: the mind reaches into regions the loud waking brain normally silos off (F19), the quiescent brain relaxing the siloing **in life** — which is the structural place to locate what the corpus calls **conversations with the dear departed**, an in-life window onto the normally-walled field. (That window is recorded here only as Iconism's mechanism; its ontology is Spirit's.) The agency-clue tells the two apart from inside: low agency reads as something happening **to** the dreamer, higher agency as something the dreamer is **doing**.

This places dreaming on an explicit ***agency axis*** with two

neighbours, indexed by where the agency for socket-shaping sits.

Waking is the normal case: downward causation operates, the self shapes its own sockets. *Dreaming* is downward causation **interrupted** — the self's agency over socket-shaping is reduced (brain-echo) or redirected (mind-origin). *Hypnosis*, new to this document, is the third point: downward causation **handed over** — the self's socket-shaping is taken up by an external arbiter, the hypnotist, who shapes the subject's sockets from outside. Waking / dreaming / hypnosis is one axis: agency with the self, agency suspended, agency with another.

Deep sleep is a configuration in which neither the sensory input nor the spirit-facts contribute strongly enough to shape sockets; the awareness-centre is not actively recentring, the handshake mechanism is largely quiescent, the icon-configuration is minimal. In the framework's vocabulary, this is the non-determined state of the handshake mechanism — no socket, no candidate, no handshake under way at the awareness-level.

That deep sleep is felt as "no consciousness" rather than as "continuous consciousness with no content" is consistent with the framework: consciousness is the obtaining of icons, and when the icon-formation mechanism is quiescent there is no consciousness to register. The waking from deep sleep is the resumption of the socket-handshake cycle, with the new sockets re-establishing the icon-configuration that the prior spirit-facts (now augmented by whatever has accumulated during sleep) bias toward continuity

with the pre-sleep configuration.

This is what the framework can say about the apparent discontinuities of consciousness without abandoning the continuity account of §7.2 and §7.3. The continuity is structural — wave-function-contribution from accumulating spirit-facts biasing the continuous Arc-selection process toward continuous icon-configurations — and the discontinuities are configurations in which the active socket-handshake cycle is quiescent. The continuity is preserved across the discontinuities by the persistence of the spirit-facts; the resumption is the spirit-facts re-shaping sockets when the brain is again in a state to produce candidate firing patterns under their constraint.

7.5 Mental particles and the question of mechanism

The Diary 064A (30 April 2026) *mental-particles / mental-force* speculation is the corpus's first attempt at a *mechanism* for how icons interact across the higher dimensions they participate in. The speculation extends the icon framework into a force-carrier picture: abstract carriers of mental force paralleling photons (which carry electromagnetic force) and gluons (which carry strong force), but in higher-dimensional space rather than in three-dimensional space.

The speculation is candidate substrate for a future paper passage on icon-icon interaction. The substrate document does not commit to the force-carrier picture, for two reasons. First,

the framework's account of icon-interaction does not require specific carriers: the wave-function-contribution mechanism of §5 supplies the structural account of how layer-3 facts interact with each other and with layer-2 activity, and that account is sufficient for the substrate's spine. Second, the force-carrier picture raises the interaction-question (how do mental-force carriers physically interact with brains?) in a form the substrate document is not ready to answer cleanly. The Princess Elizabeth question, in this section's context, is laid in the form §8 develops, which does not depend on a force-carrier picture.

What the substrate document does record is that the mental-particles speculation, if developed, would extend the framework's account of icon-icon interaction beyond the wave-function-contribution mechanism. The wave-function-contribution gives the structural account of how layer-3 facts bias layer-2 outcomes; mental-force carriers, if they exist in the framework's sense, would give the layer-3-to-layer-3 mechanism by which icons interact with each other directly, without going through layer-2. This is not currently required for any working document in the chain, but it is candidate substrate for future development.

The flag is that the speculation deserves engagement at some later stage. The substrate document neither commits to it nor rejects it; it positions it as a future-work item that the substrate's existing apparatus does not require.

7.6 Status of the section's claims

The section's claim is that continuity of consciousness across Arcs is **structurally accounted for** by the wave-function-contribution mechanism operating continuously, the centre-of-awareness recentring without discontinuous jumping, and the icons persisting across Arc-by-Arc layer-2 substrate changes. This is the framework's positive sketch.

Whether the sketch constitutes a **complete** answer to the second hard problem is a question the substrate document leaves open. The sketch shows that the framework has structural resources to account for continuity; whether those resources are sufficient to capture the full phenomenology of conscious persistence is something ``centre_of_awareness.md`` and ``consciousness_awareness_spirit.md`` will be in a better position to assess once they are developed.

The section also leaves open the question of what happens to consciousness at death (and in any analogous "in-between" states the corpus has flagged, e.g. the 11 May 2026 **Spirit, DC, and the in-between state** passage of Diary 064B). The framework's position on death — that what survives is the spirit-facts as eternal facts, with the conscious-self in its real-and-virtual form having demerged — is Spirit-side material, not Iconism-side material, and the substrate document leaves it to the eventual Spirit project to develop.

What the substrate document needs is a positive treatment of continuity in the layered-emergence vocabulary, which is what §7.2 and §7.3 have supplied. The sketch is the framework's current position; the working documents that follow may refine it, but they should not start from scratch.

8. Princess Elizabeth answered

Princess Elizabeth of Bohemia, in her letters to Descartes in 1643, pressed the question that Cartesian dualism has not in three and a half centuries been able to answer cleanly: how, on the dualist's account, can the immaterial mind move the material body? If the two are different in substance, by what mechanism does the one act on the other? Descartes's appeal to the pineal gland was a gesture rather than an answer; subsequent dualist accounts have tended either to repeat the gesture in different anatomical locations or to retreat into denial that the interaction problem is well-posed.

The framework's position is that the interaction problem is well-posed for Cartesian dualism and is fatal to it, but is not the correct shape of question to ask of Iconism, which is not a dualism. The substrate document is in a position, by §§2–7, to state the framework's response to Princess Elizabeth precisely. This section does so.

The corpus has been working toward this answer since Diary 036 (April 2019). The 21 December 2020 passage of Diary 042 was the first sustained attempt; the 22 October 2021 *Downward Causation Inset Day* of Diary 046 supplied the photon-probability mechanism; Diary 049 (August 2022) restated the answer with icons and Sperry's Waterwheel; Diary 059A (January 2025) sharpened the wave-function-guess-and-confirm picture; Diary 060A (April 2025) gave the *memory is spirit is facts* formulation; Diary 061 (19 September 2025) gave the explicit *Answering Princess Elizabeth* passage with *the brain is dumb plumbing* as its candidate aphorism; Diary 064B (12 May 2026) gave the handshake account; the addendum (26 May 2026) gave the HRM identity statement. The substrate document's task is to state the answer in operational form, drawing all these strands together.

8.1 The question restated

Princess Elizabeth's question presupposes Cartesian dualism: two substances, the mental and the physical, fundamentally distinct, each with its own kind of properties. If they are this distinct, how can either causally affect the other? The interaction would require a bridge between substances of fundamentally different kinds; no such bridge has ever been described that does not either (a) collapse one substance into the other or (b) introduce a third bridging-substance which itself raises the same problem.

The framework's response is structural: the question is not well-posed for an account that is not a substance dualism. Iconism is not a substance dualism. There is one substance — Virtualism's relationships-between-virtual-wholes

(`foundations\_v001.md` §6) — and three emergent layers of fact about that substrate (§2 above). Mind is not a different substance from brain; mind is what the brain-with-its-firing-patterns *is* at the layer where indiscernibility-with-worldly-wholes can hold. The interaction problem in Princess Elizabeth's form requires two substances to interact across; the framework has only one substance to begin with.

This is the structural shape of the answer. The sections of the substrate document have been preparing it.

8.2 One substance, three layers

§2 set out the three-layer picture. The material brain (layer 1), the brain's activity (layer 2), the icons and mind (layer 3) are not three substances. They are three levels of fact about one substrate, related by Virtualism's emergence apparatus. Layer 2 supervenes on layer 1; layer 3 supervenes on layer 2. Each layer carries real facts, but the facts at each layer are facts *about* the substrate, not facts in a separate substance.

The work of the layered-emergence picture is to refuse the either/or that the standard interaction problem presupposes. The mind is not a separate substance (so it does not have to find a

way to interact with the brain across a substance-divide); but the mind is also not nothing-over-and-above the brain (so it can have its own causal contribution). The mind is the brain-at-the-layer-of-icon-formation. The interaction is not between substances; it is between layers of fact about one substrate.

This is what the §2 discipline buys. The three layers are not just an organisational convenience; they are the framework's structural commitment about what exists in the brain-mind system. There is one substrate. There are three layers of fact about it. There is no second substance for the mind to be made of, and therefore no interaction-across-substances for the framework to have to explain.

8.3 Layered emergence is not reduction

A natural objection at this point: if the mind is no separate substance, is it not just identical with the brain after all? Have we not collapsed back into reductive physicalism, with all the problems that view has trying to do justice to the phenomenology of consciousness?

The answer is no, because layered emergence is not reduction. The icon at layer 3 is not identical with any firing pattern at layer 2. The same icon can be supported by different firing patterns (multiple realisation, preserved by §2.5 above); the same firing pattern would not, in a different brain or in a different context, produce the same icon (because the icon is

individuated by its layer-3 relations, which include relations with other icons that may not be present in a different brain). The layers are related by emergence; they are not related by identity.

This is the structural shape `paradox\_and\_emergence\_v003.md` §5 develops: emergence in the framework's strong sense is a real relationship between layers, with the higher layer carrying facts the lower layer does not carry on its own. The higher layer is not a separate substance — it is layers of fact about the same substrate — but the higher layer's facts are not reducible to the lower layer's facts taken individually. The layered emergence picture is the third option that reductionism and substance-dualism between them fail to make available.

This third option is what permits the framework to give a positive account of psychological causation without either collapsing mind into brain (which would deny the mind's causal contribution) or separating mind from brain (which would re-pose Princess Elizabeth's problem). The mind is *at a different layer* from the brain; it has its own dynamics, its own facts, its own causal contributions; but it is the same substrate, layered.

8.4 Downward causation as wave-function constraint, not interaction

§5 developed the framework's mechanism for downward causation: the layer-3 facts of the mind contribute to the wave functions

that the brain's quantum events collapse under. This is not interaction between substances. It is constraint on a single wave function by facts at different layers of organisation within a single substrate.

The wave function in the brain's vicinity is what it is — a structured possibility-space for the brain's imminent future. What determines that structure is *all the relevant facts in the relevant region*: the layer-1 facts about neurons and synapses, the layer-2 facts about firing patterns currently underway, the layer-3 facts about icons currently obtaining, the spirit-facts (in their wave-function-contributing form) about the prior history of the individual. All these facts contribute to the wave function. None is in competition with any other; they jointly determine what the wave function is.

When the wave function collapses at an Arc, the outcome is one of the possibilities the wave function supports. The outcome is biased by the joint contribution of all the facts. The layer-3 facts do not "intervene" on the wave function from outside; they are part of what the wave function *is*. The collapse is not the mind acting on the brain; it is the universe selecting one outcome from the wave function that the brain (in its full layer-1-through-layer-3 specification) is supporting.

This is the wave-function-constraint mechanism stated as the response to Princess Elizabeth. There is no interaction across

substances because there are no two substances. There is no interaction across layers because the layers are not separate substrates that have to communicate; they are layers of fact about one substrate. The wave function is the substrate's local imminent-future-structure; the facts at all three layers contribute to that structure; the collapse selects an outcome biased by the joint contribution. The mind influences the brain because the facts of the mind are among the facts the wave function answers to.

8.5 The brain as parts-supplier; the mind as socket-shaper

§4 developed the handshake mechanism: the mind creates a socket at layer 3 (an under-determined shape for an icon to fit); the brain produces firing patterns at layer 2 under wave-function-constraint; an icon at layer 3 emerges when a firing pattern's icon fits the socket. This is the operational form of how mind and brain co-operate without interacting-across-substances.

The 12 May 2026 *consciousness-as-handshake* passage of Diary 064B states the picture in the corpus's own vocabulary: *mind as socket, brain as parts-supplier*. The 19 September 2025 *Answering Princess Elizabeth* passage of Diary 61 states it more pungently: *the brain is dumb plumbing*. Both formulations capture the same structural fact: the brain does its own work under its own (material) constraints; the mind does its own work of socket-shaping; the join between them is the wave-function-contribution that biases the brain's work toward producing

icon-supporting patterns. The mind does not push the brain around; the mind shapes the constraint that the brain operates under.

This is the answer in operational form. The brain is doing what brains do — material processes governed by the brain's own mechanics. The mind is doing what minds do — shaping sockets at layer 3, recentring on completed icons. The wave-function-contribution is the join: it is how the layer-3 facts (the mind's facts, the spirit's facts) bias the layer-2 outcomes (the brain's firing patterns) without requiring the mind to *act across a substance-divide* on the brain. The mind's facts are in the same substrate as the brain's facts; they contribute to the same wave function; the bias is the contribution.

8.6 The same gravity process in unlimited dimensions

The framing claim of §1 — that Iconism is Virtualism applied to the brain-icon case, using the same machinery, in many more dimensions, with both attraction-by-sameness and repulsion-by-difference operating in each dimension — applies precisely to the Princess-Elizabeth-answering case. Gravity is the structural parallel; the centre of gravity of a system of masses constrains the parts of that system by being where it is, and the parts respond to where it is. There is no "interaction" of the centre with the parts in the sense Princess Elizabeth's question requires; the centre is a virtual fact about the configuration, the parts are arranged to be consistent with the

centre's being where the configuration makes it be.

The same is true of the icon-and-mind case at layer 3. The awareness-centre, like the centre of gravity, is a virtual fact about the configuration of icons. The icons, like the masses, are arranged to be consistent with the centre's being where the configuration makes it be. The brain's activity is biased toward producing firing patterns whose icons are consistent with the existing configuration, not because the centre is *acting on* the brain, but because the centre is a fact about the configuration that the wave-function-contribution mechanism incorporates.

The 26 May 2026 addendum's HRM identity statement is the strongest single corpus formulation: *the brain's reasoning is the inside-the-brain gravity process, exactly the same process as physical gravity but in unlimited dimensions*. Princess Elizabeth's question, on this reading, is the question one would ask of physical gravity if one took the centre of gravity to be a separate substance acting on the masses: how does the centre of gravity move the masses? The answer is that the centre does not "move" the masses; the masses are arranged in the configuration that has the centre where it is, and as the masses move the centre moves with them, with the centre's location being just *where the configuration of masses puts it be*. There is no interaction-question to be asked because there are no two substances.

This is what laying Princess Elizabeth for good amounts to. Her question was the right question to put to substance dualism. It has no answer for substance dualism, which is why substance dualism is wrong. But Iconism is not substance dualism. The question, asked of Iconism, is asking about an interaction that does not occur because there are no two substances to interact. The question dissolves; what is left is the structural mechanism the framework gives — layered emergence, wave-function constraint, the gravity process in unlimited dimensions — and that mechanism gives the operational account of how mind and brain co-operate without interacting-across-substances.

8.7 What this section has closed

This section has closed Princess Elizabeth's question by stating the framework's structural position: one substance (Virtualism's relationships-between-virtual-wholes), three emergent layers of fact about it (material brain, firing patterns, icons-and-mind), downward causation by wave-function-contribution from higher-layer facts on lower-layer outcomes, the icon-mind layer as the gravity-process-in-unlimited-dimensions. There is no second substance for the mind to be made of; there is no interaction-across-substances for the framework to have to explain; there is only one substrate, layered, with facts at every layer contributing to the wave functions the universe's Arc-selection process runs on.

The substrate document closes here. The remaining working documents — `icon\_and\_world.md`, `consciousness\_awareness\_spirit.md`, `centre\_of\_awareness.md`, `conceivability.md`, `ai\_and\_sentience.md` — will develop their respective topics using the apparatus this document has supplied. Where they require the layered-emergence picture, they refer to §2 here; where they require the handshake mechanism, to §4; where they require the wave-function-constraint account, to §5; where they require the icon-as-high-dimensional-star geometry, to §3; where they require the continuity-across-Arcs sketch, to §7; where they require the framework's position on Princess Elizabeth, to this section.

The substrate is in place. The Iconism working-document chain can proceed.

9. What this document does not do

The substrate document supplies the apparatus the working-document chain needs. It does not do the work of the chain's subsequent documents, and several topics are deliberately deferred to where they belong.

9.1 Topics deferred to the working-document chain

- **Icon and world.** The relationship of icons to the world they represent — when an icon is *of* a thing, what the correctness

conditions on representation are, how icons can be of things that do not exist, how icons can misrepresent. Goes to `icon\_and\_world.md`. The substrate document gives the high-dimensional-star geometry and the wave-function-constraint account of how icons exert effect; what an icon *is of*, and how reference works, is the world-side story that the next working document develops. (Some passages in §3 above touched on this — the NATO-star convergence, the indiscernibility-up-to-the-feature-set — but the full account of representation is not here.)

- **Consciousness, awareness, spirit.** The three-term distinction that the *Philosophical Log* draws — consciousness as the capacity for awareness, awareness as the actual having of experience, spirit as the standing pattern of facts (memories, dispositions, characteristic styles of attention) — is used freely in this document but is not developed. The substrate document treats these as picked up from the Log; the working document `consciousness\_awareness\_spirit.md` will give the full treatment, including the questions about whether consciousness is graded or threshold, whether awareness has degrees, and what exactly spirit's relationship to fact-causation is.

- **The centre of awareness.** What the centre of awareness *is*, in framework-internal terms, and how it relates to the centre-of-gravity machinery from `gravity\_and\_the\_past\_v1\_0.md` §13.6. This document has used the centre-of-awareness phrase in

§4.5 and §7 without arguing for the analysis; the analysis is the working document `centre\_of\_awareness.md`. The expected shape of the analysis — centre-of-awareness as virtual fact about an awareness-configuration, just as centre-of-gravity is virtual fact about a mass-configuration — is what makes the continuity-across-Arcs sketch in §7 work, but the document for that analysis is downstream.

- **Conceivability.** What it is for something to be conceivable, how conceivability relates to possibility, what the Two-Door argument's exploitation of conceivability comes to, and how the framework's account of icons supplies the right reading of "conceivable but not actual". Goes to `conceivability.md`. The Iconism paper (downstream of the working-document chain) is expected to use this material centrally.

- **AI and sentience.** Whether AI systems can be conscious, and what the framework's analogue-requirement claim implies for silicon implementations versus biological ones. The substrate document has been silent on this — deliberately, to keep substrate from doing working-document work. The framework's position (analogue substrate required, digital simulation insufficient) is set out in the *Philosophical Log* and is to be developed in `ai\_and\_sentience.md`.

9.2 Topics deferred to the Tier 1 documents

- **What virtual facts are.** Goes to `foundations\_v001.md` §6.

The substrate document has taken Virtualism's metaphysics as given and used it to frame the layered-emergence picture; it has not argued for the metaphysics.

- **Four-state determination apparatus.** Goes to ``paradox_and_emergence_v003.md`` §3. The substrate document has applied the apparatus to the handshake mechanism in §4 but has not argued for the apparatus.

- **Wave functions as application of facts.** Goes to ``change_and_present_v1_0.md`` §4. The substrate document has used the wave-function picture in §4 and §5 to give the downward-causation account and the guess-and-confirm account of recall; it has not argued for that picture's general role in Arc-selection.

- **Arcs and the present.** Goes to ``change_and_present_v1_0.md`` §2. The substrate document has invoked Arcs in §4.5 and §7 in giving the continuity-across-Arcs sketch; it has not argued for the Arc apparatus.

- **Strong/weak emergence and the layered ontology.** Goes to ``paradox_and_emergence_v003.md`` §5. The substrate document has used the strong-emergence framing to give the layered-emergence picture but has not argued for that framing.

9.3 Topics deferred to future work

- **Mental particles.** Flagged in §7.5 as a speculation in the diary material (Diary 064A 30 April 2026) — that the framework might require a mental analogue of force-carrying particles to give a clean account of how higher-layer facts make their wave-function contributions felt. The substrate document does not commit to this picture, and the wave-function-contribution mechanism as stated does not require it. If the working-document chain or the paper finds that some such picture is needed, it will be developed there.

- **The "in-between" state question.** The diary records (Diary 064B, 11 May 2026) a question the author has not resolved: whether, between Arcs, there is some structural sense in which the icon-mind is "between" determinate states, or whether the Arc-selection process is fast enough that no such between-state is needed. The substrate document has been silent on this. The question may be answered by the continuity-across-Arcs working out, or it may persist as an open question into the paper.

- **The empirical interface.** What experiments could in principle bear on the framework's claims. The substrate document is philosophical throughout; what neuroscience could be done that would tell for or against the layered-emergence picture, the high-dimensional-star geometry, the wave-function-constraint account, is a question the working-document chain or the paper may want to engage. The diary material gestures at this in several places (Diary 063B's Bell's-inequality-structural-

argument is the clearest case) but no programmatic engagement has been done.

- **Engagement with the literature on mental causation.** Kim's Causal Exclusion argument has been engaged structurally in §5; the full engagement with Kim, Davidson, Yablo, Jaegwon-Kim's successors, and the wider mental-causation literature has not been done. The Iconism paper is the right place for that engagement.

- **Engagement with phenomenology proper.** The substrate document has used phenomenological cases (the orange-perception case, the tip-of-the-tongue case, the recall cases) as data-points the framework should match, but it has not engaged with phenomenology as a tradition. Husserl, Merleau-Ponty, Sartre, Heidegger, and the analytic phenomenology that has followed are all candidate engagement-points for the consciousness-awareness-spirit working document or the paper.

10. Status and open questions

Status

The substrate as set out here is stable across the diary material from Diary 046 (October 2021, the Downward Causation Inset Day) to Diary 064B (May 2026) and the *Philosophical Log* v6. The apparatus has been developed, refined, and tested over

approximately four and a half years of diary work, and the formulations the substrate document uses are the mature formulations as of the 26 May 2026 addendum to Diary 064B.

The substrate's two structural claims —

1. that the icon-and-mind layer is Virtualism's gravity-process in unlimited dimensions, the same machinery the physical world uses, not an analogous machinery; and
2. that the unified mechanism by which the layer operates — the four-state-handshake of icon-as-socket with brain-as-parts-supplier, the wave-function-constraint account of downward causation, and the recentring of awareness in successive Arcs — gives a worked operational account of mind-brain co-operation that dissolves Princess Elizabeth's interaction question

— are firm. The detailed working-out in §§2-8 has been done in the diaries and is summarised faithfully here. What remains is the working-document chain's development of the topics §9.1 defers, and the paper's positioning of the framework in the mental-causation literature.

A reasonable summary paragraph for the command project's ``theory_current_state_v004.md`` (per the directive note in `v002_1`):

- > *Iconism, the framework's working-out of consciousness in
- > Virtualism's terms, claims that the icon-and-mind layer is the
- > gravity-process the framework uses for physical mass and
- > energy, run in unlimited dimensions on the firing-pattern
- > substrate the brain supplies. Three emergent layers — material
- > brain, firing patterns, icons-and-mind — co-operate via a
- > four-state handshake in which the mind shapes an icon-socket,
- > the brain delivers candidate firing-pattern parts, and the
- > Arc-selection process determines which candidate fills the
- > socket. Downward causation operates as wave-function-
- > contribution: facts at the icon layer contribute to the wave
- > functions the universe runs on, with no exclusion of physical
- > causes at the firing-pattern layer because both layers
- > contribute alike. The continuity of consciousness across Arcs
- > is given by the structural-persistence of the icon-and-spirit
- > configuration that defines a centre of awareness, in the same
- > way the structural persistence of a mass-configuration defines
- > a continuous trajectory of a centre of gravity. Princess
- > Elizabeth's interaction question dissolves: there are not two
- > substances, there is one substrate layered into facts at three
- > emergent levels, and the question of "how does the mental
- > interact with the physical" is asking about an interaction
- > that does not occur because there are no two substances to
- > interact.*

Open questions

- **The continuity-across-Arcs analysis is a positive sketch,

not a final resolution.** §7 supplies the framework's account of why successive Arcs do not produce phenomenal discontinuity, but the account leans on the centre-of-awareness analysis that the working document `centre\_of\_awareness.md` has yet to give. The sketch is firm enough that the substrate-document can reasonably present it; the full account waits on the working document.

- **The in-between state question.** Flagged in §9.3. The question is whether the Arc-selection process is fast enough that no between-state of the icon-mind is needed, or whether the framework should countenance some structural between-state. The diary has not resolved this.

- **The mental-particles speculation.** Flagged in §7.5 and §9.3. The diary records the speculation that some mental analogue of force-carrying particles might be needed to give a clean account of how facts at the icon-mind layer make their wave-function contributions felt. The substrate document does not commit to this picture and does not require it. Whether the framework will eventually want it is open.

- **The supervenience-without-exclusion claim.** §5.4 has completed the *Hard Problem* essay's Supervenience sub-section by giving the framework's positive answer to Kim's Causal Exclusion challenge: layer-3 facts supervene on layer-1 facts but are not exhausted by them, and supply their own

wave-function contributions over and above the contributions layer-1 facts make. The structural claim is firm; what remains open is the full engagement with the supervenience literature (Kim's later work, the multiple-realizability tradition, Yablo's proportionality account) which is paper-level work.

- **The Bell's-inequality-structural-argument.** §5.5 has given the structural argument that Kim's exclusion principle is in the same family as the hidden-variables programme that Bell's inequalities ruled out: it insists on a kind of completeness at one layer that the framework's wave-function-contribution account denies. The argument is suggestive rather than decisive in its current form; the paper-level engagement with Bell's-inequality literature and with the structural-realist tradition that has used Bell's-inequality arguments in philosophy of mind has not been done.

- **The empirical interface.** Flagged in §9.3. What experiments could in principle bear on the framework's claims is a question the substrate document has not engaged. The diary's gestures are suggestive (the Bell's-inequality move, the Sperry-wheel illustration of emergent rotation, the three-photons-three-neurons triangle) but no programmatic engagement has been done.

Candidate aphorisms

For the bank, from the diary material this document has drawn on:

- "A balanced mind can remember easily." (Diary 064B, 8 May 2026.) Picks out the well-tuned-handshake claim in one line. The "easily" is doing significant work: a mind that has its icon-shaping in good order does not have to **force** the handshake by clenching attention or running compensatory searches. The trying-too-hard pathology that §4.3 and §6.2 give the account of is precisely the case where ease is absent.

- "The brain is dumb plumbing." (Diary 61, 19 September 2025, **Answering Princess Elizabeth**.) The brain-as-parts-supplier claim in one line. The "dumb" is not dismissive: the brain is a remarkable piece of plumbing, doing the work the framework needs done at the firing-pattern layer. What it is not is the seat of mind. The mind is the icon-layer configuration, of which the brain's firing-pattern activity is the parts.

- "The mind creates a socket for the brain to fill or fuck over." (Diary 064B, 12 May 2026.) The full handshake claim, including its failure modes — the brain can fail to deliver parts, can deliver the wrong parts, can deliver parts that fit poorly. The phrase's vulgarity is preserved because the diary insists on it and because the substantive point is the failure-mode-inclusion: an idealised handshake-always-succeeds story would not match the phenomenology.

- "The brain's reasoning is the inside-the-brain gravity process, exactly the same process as physical gravity but in unlimited dimensions." (Addendum to Diary 064B, 26 May 2026.)
The framing claim — Iconism as Virtualism applied — in the diary's own words. This is the candidate centrepiece formulation for the Iconism paper.

- "Memory is spirit is facts." (Diary 060A, 8 April 2025.) The three-term identity that underwrites both the continuity-across-Arcs sketch and the Alzheimer's-as-loss-of-retrieval account. Memory is not stuff stored in the brain; memory is the standing pattern of facts that the brain has been shaped by, and that the brain's firing-pattern activity at any given moment may or may not retrieve.

- "Downward causation works by altering the chances of photons going one way rather than another." (Diary 046, 22 October 2021, paraphrased — the *Downward Causation Inset Day*.) The framework's foundational claim for downward causation, given six months before the *Philosophical Log* was written. Icons as virtual obstacles, altering photon-probabilities, is the pre-formulation of the wave-function-contribution account §5 has developed.

- "What it is to *be* the icon, the seeing of orange, is what it is for a mind to inhabit the firing pattern in question

with its full structural content." (Diary 064B 12 May 2026, paraphrased.) The framework's positive answer to the hard problem of consciousness in inhabitation-terms rather than identity-terms: the firing pattern and the icon are not two things to be identified; the icon is what the firing pattern *is*, at the layer the icon is identified at.

11. Cross-references summary

Where each Tier 1 document is used:

- *foundations_v001.md* — §6 (everything reduces to virtual facts) is the substrate of §1's framing claim. §5 (Leibniz's Law) is the substrate of §3's selective-indiscernibility treatment of icon identity. §7 (categorical distinctions) is in the background throughout: icons are *facts*, of which there are different kinds (true, factual, real, ideal, virtual), and the icon-mind layer is the place where the categorical vocabulary is doing most of its work.

- *paradox_and_emergence_v003.md* — §3 (the four-state determination apparatus) is the substrate of §4. §5 (strong/weak emergence) is the substrate of §2's layered-emergence picture, with the icon-mind layer being a case of strong emergence as the document analyses it.

- *gravity_and_the_past_v1_0.md* — §13.6 (centre of gravity

as virtual fact about a mass-configuration) is the substrate of §7's continuity-across-Arcs sketch, with centre-of-awareness standing to icon-configuration as centre-of-gravity stands to mass-configuration. The "gravity process in unlimited dimensions" framing claim depends on the gravity process being well-developed elsewhere, which it is in this document.

- `change_and_present_v1_0.md` — §2 (the Arc apparatus) is the substrate of §4.5 and §7's continuity-across-Arcs sketch. §4 (wave function as application of facts) is the substrate of §4's guess-and-confirm account of recall and §5's wave-function-contribution account of downward causation.

Where each working-document chain document is anticipated:

- `icon_and_world.md` — picks up the representation question that §3.4 and §3.5 have gestured at without resolving. The high-dimensional-star geometry the substrate document has developed is the apparatus the world-side story will use.

- `consciousness_awareness_spirit.md` — picks up the three-term distinction used freely throughout, especially in §7's continuity-across-Arcs sketch and §4.5's recentring-of-awareness account. The substrate document has assumed the distinction; the working document will develop it.

- **centre_of_awareness.md** — picks up the centre-of-awareness machinery the substrate document has used in §4.5 and §7 without analysing. The analysis — centre-of-awareness as virtual fact about an awareness-configuration, by analogy with centre-of-gravity as virtual fact about a mass-configuration — is anticipated but not given here.

- **conceivability.md** — the substrate document has not engaged with conceivability directly, but the icon-as-high-dimensional-star apparatus is the substrate the conceivability working document will use. Conceivability is, in the framework's vocabulary, a fact about which icon-configurations can be assembled at the icon-mind layer without internal contradiction.

- **ai_and_sentience.md** — picks up the silicon-versus-carbon question the substrate document has deliberately not engaged. The framework's analogue-requirement claim (consciousness requires the right kind of substrate, not merely the right pattern of information-processing) is developed in the *Philosophical Log* and will be the working document's central topic.

Where the *Philosophical Log* v6 is used directly:

- Icon definition (used in §3).
- Consciousness / awareness / spirit definitions (used in

§4.5, §7; deferred to `consciousness\_awareness\_spirit.md` for full treatment).

- The p-figures discussion (used in §3 implicitly; the high-dimensional-star geometry is the framework's mature form of what the Log called p-figures).
- The AI-and-analogue-requirement passage (deferred to `ai\_and\_sentience.md`).
- The Spirit-and-fact-causation passage (used in §5's wave-function-contribution account, where Spirit's role is presupposed but not developed; deferred to `consciousness\_awareness\_spirit.md`).

12. Change-log: v001 → v002

This pass folded the ratified refinement findings into the document in a single coordinated edit. Each entry below names the change, the section it landed in, and the finding or decision that drove it.

Findings notes: refinement v002–v004 (F1–F23), v005 (D10), v006 (D5, F28–F29); decisions D6 (dreams), D9 (core claim), D10 (determination traversal), D5 (shape split).

1. **§3.1 — pointer principle.** Stated that every icon-dimension is a generated pointer the icon instantiates none of; the mass-dimension is a created pointer, not a mass in the head. Reinforced in §3.3 (the star's points are pointers at what the object instantiates). *Drivers: F1.*
2. **§3.1 — shared/created split.** Distinguished the shared,

- anchoring dimensions (position, shape) from the created, added dimensions (colour), with the tier-placement reason and the wholes/parts gloss for shape. *Drivers: F1, F2, F28, D5.*
3. **§3.1 — core claim.** Stated the D9 core-claim sentence verbatim, with the F22 "worldly = primal" gloss in the surrounding text.
Drivers: D9, F22.
4. **§3.1 / §4.3 — the determination traversal.** Presented the four states as motion along the scale, added the synchronic anatomy and the delineated/undelineated void split, reworded "fully-determined = at rest" as *just-settled*, and added the perception traversal (under-determined rearrangement → closure → lit settled fact) as the worked example. *Drivers: D10, F10, F11; F19/F20 (closure), F24 (lighting) cross-referenced.*
5. **§5 (intro, §5.1) — downward-causation nuances.** Added the F6 perspectival nuance (up/down as one constraint), named Princess Elisabeth alongside Kim as one challenge, and added the upward half of the virtuous circle. *Drivers: F6, F8, F9.*
6. **§6.3 — recall riders.** Added settling-on-apparent-fit and the handshake-as-local-Arc reading, with the relaxed-self corollary.
Drivers: F7, F12.
7. **§3.1 / §4.3 — F2/F3/F10 material.** Folded the Venn (shared/created), the shared-ontology grounding, and the determination-scale anatomy inline rather than as a separate sub-section. *Drivers: F2, F3, F10.*
8. **§6.3 — the headline rewrite.** Split recall into the forceless Leibniz connection (the answer-shaped void filling itself, no carrier, distance-independent) and the brain's guess-and-paint

- step (where wave-function language stays) — the fix for the
"shouting wave function." *Drivers: F13.*
9. **§4.6 — the graded family.** Re-indexed the cognitive family by
question/answer icon-counts, added the *sense* entry, named will as
the shared element of the upper members, and added the
learned/remembered/imagined fact-taxonomy with fictions as
partially-untrue facts. *Drivers: F14.*
10. **§4.5 — the Cartesian Theatre reframed.** Added the Dennett
positioning (no place, a real abstract configuration; centring =
brightness by meaningful connection), tied to the
centre-of-gravity-as-virtual-fact move. *Drivers: F17.*
11. **§7.4 — dreams and hypnosis.** Split dreams by origin/agency
(brain-echo vs mind-origin), added hypnosis as the
external-arbiter case, and stated the waking/dreaming/hypnosis
agency axis. *Drivers: D6, F15, F16.*
12. **§2.1 — constraint ladder and mind/spirit rider.** Added the
upper half of the constraint ladder (activity → icons →
consciousness → awareness) as the upward extension of the
three-layer emergence, and the F18 same-kind/different-extension
rider where consciousness/awareness/spirit is first stated.
Drivers: F18, F20.
13. **§6.3 — the siloing mechanism.** Stated that the brain's
brightness silos the mind, with the post-mortem un-siloing
consequence flagged (not developed) for Spirit. *Drivers: F19.*
14. **§4.5 — awareness maintains and replicates.** Added that
awareness replicates its icons (keeping them recent, bright,
close), the reason rehearsal aids recall; awareness as both

wholeness and heart of mind. *Drivers: F21.*

15. **§5 (intro) — downward causation as wholes-on-parts.** Stated DC

as the action of any whole on its parts on a stability gradient,
none absolutely permanent (the proton may decay), with the
stable-not-eternal register held apart from Spirit's eternal-fact
register. *Drivers: F23.*

16. **§7.4 — all dreams are experiences of awareness.** Added the

through-line that awareness is continuous through sleep, only
origin and agency varying. *Drivers: F16, F21.*

Open items not addressed by this pass (carried): O2 (what
intrinsic-overdetermination pressure does to a boost-less icon;
Iconism, non-blocking); O7 (the constraint hierarchy as a Tier-1
claim) and O8 (force carriers and the four I's) — both carried to the
Virtualism update. The later F24/F25 downward-causation pass (icon
intensity as extrinsic brightness; no microtubules required) and the
F26/F27 material (the AI criterion; width/length) are downstream of
this pass and feed `the\_challenges.md` and `ai\_and\_sentience.md`
rather than the substrate document.

*The substrate is in place. The Iconism working-document chain
proceeds from here.*