

Table of Contents

Icon and World.....	1
0. Why this document exists	2
1. The core claim.....	2
2. Leibniz's Law in the author's version	3
3. What the icon is indiscernible <i>from</i>	4
4. The geometry: star, hollow, and boundary.....	6
5. What the brain adds	8
6. Reference, and why this is not representationalism	9
7. Error, fiction, and the limits of anchoring.....	10
8. Multiple realisation.....	12
9. The Cartesian theatre, reframed.....	12
10. What this document does not do	13
11. Status and open questions.....	13
12. Candidate aphorisms.....	14

Icon and World

Working document, Iconism Theory project. The second of the working documents to stand on icon_dynamics_v002.md, and the one that carries the framework's core claim. Sets out what an icon is indiscernible from and in what sense; states Leibniz's Law in the author's version and names the idiosyncrasy rather than softening it; resolves the star and the hologram into one structure read from two ends; gives the hollow its measure; says what the brain adds over a mirror; and answers the standing objections to any account that makes likeness the basis of aboutness.

Sources: D9 and icon_dynamics_v002.md §3 for the geometry and the core claim; the author's note On the nature of Icons (21 August 2026) for the star-hologram resolution and the matching-of-internality formulation; F52, F55, F56, F57, F58, F59 (dialogue findings, 21 August) for the corrections and refinements the note produced; F22, F28, F38, F46, D5, D10 for the shared/created split, reference and fiction; Iconism Interface v002 §1, §4, §5 for what Tier 1 supplies. Where a formulation is Claude's inference rather than the author's statement it is marked as such.

Nothing here has been written back into icon_dynamics_v002.md. §4.3, §6.3 and §7.3 contain proposals awaiting the author's ruling.

0. Why this document exists

`icon_dynamics_v002.md` establishes how icons form, decay and act. It presupposes throughout that an icon can stand in an indiscernibility relation with a worldly whole, and defers the account of that relation to here.

The deferral was deliberate and the debt is large. Indiscernibility is the framework's central claim, it is what distinguishes Iconism from every rival in the field, and it is also the claim most exposed to a family of objections that have sunk resemblance theories of mind since Berkeley. A version of this document that stated the claim without answering those objections would be the framework's weakest point rather than its strongest.

So this document does four things. It states the relation precisely, including what the icon is indiscernible *from*, which is not what the earlier formulations said. It distinguishes the two relations an icon actually holds — likeness on some of its dimensions, pointing on others — because running them together is what makes the account look naive. It gives the geometry its measure, including the mathematics that turns the star from an image into a result. And it answers Berkeley, Goodman, Putnam, Fodor, Kim and Dennett, in that order of difficulty.

1. The core claim

1.1 D9, stated

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.

The author's own compact gloss, given to a general audience in August 2026 and better than any paraphrase of it: *"I use icon in the religious sense, i.e. that the image is so like the reality that the two have the same factual properties."*

That sentence should open the paper's treatment. It says what the word is doing — which a referee will ask — and it puts the claim at the strength the framework actually holds it. Not that the image resembles the reality. That the two have the same properties.

1.2 Two relations, not one

This distinction is implicit in `icon_dynamics_v002.md` §3.1 and is stated explicitly here because running the two relations together is the single easiest way to misread the theory.

An icon does not hold one relation to its object. It holds two, on different dimensions, and they are different in kind.

On the shared dimensions the relation is likeness — strictly, indiscernibility.

There is a real Euclidean geometry out there with prior existence. The icon's position-fact and shape-fact are indiscernible from it, and by Leibniz's Law indiscernible features are *the same feature*. This is an identity claim about features, not a similarity claim about objects.

On the created dimensions the relation is pointing. Colour is constructed from the proportions in which three cone types are triggered; magenta has no wavelength, so colour is not in the light and there is no Tier 1 fact of colour's own kind for the icon to be indiscernible from. The colour-dimension is a *handle* onto a real optical property the organism could not otherwise register. It is about something without being like it.

The consequence for the paper is large. Almost every classical objection to likeness accounts assumes that the account needs likeness to do *all* the work — to supply aboutness across the board. It does not. Likeness anchors; pointing extends. **The icon is like its object where it shares dimensions with it, and about it where it does not,** and the two are held together by being dimensions of one whole.

This also states, in a sentence, what Iconism has that Integrated Information Theory does not: a place to put aboutness at all.

1.3 What “worldly” means

The world-anchored perceptual case is **primal**, not restrictive (F22). It is the evolutionary origin of icons and the case in which the machinery is easiest to see. Icons extend from it to the self, to the past, to deduction's conclusions and to invented worlds. Along that extension the anchor weakens and the created component dominates, with fiction as the limit case — an icon whose worldly anchor has gone to nothing.

“Worldly object” therefore names the primal anchored case. It is not a claim that icons are only ever of external physical things.

2. Leibniz's Law in the author's version

2.1 Features, not objects

The literature reads the Identity of Indiscernibles as a principle about whole objects: if *a* and *b* share all properties, $a = b$. The framework applies it to **features**. Where two things are indiscernible in some feature, that feature is the same feature in both — *not two similar greens but one green*.

The author has flagged that this use is idiosyncratic and adopted “*out of respect as the meaning is so close.*”

2.2 Why the idiosyncrasy should be named rather than softened

Claude's recommendation, made in the 21 August dialogue and recorded here because it bears on how the paper is written.

It costs less than it appears to. Applying indiscernibility to features rather than to whole objects, and holding that a shared feature is *one* feature and not two similar ones, is a **realism about immanent universals** — a position with a substantial literature behind

it, and one on which the Law then applies perfectly binarily. Two wholes are not identical; the feature they share is. Nothing is fudged, and the framework should say which position it holds rather than let a referee decide it is being loose with a famous principle.

The framework also has more standing here than it thinks. Tier 1 already argues over Black's two iron balls, which is the identity-of-indiscernibles debate proper.

2.3 The genuine extension, and it should be named separately

The part that goes beyond immanent-universals realism is this: **the shared feature is a true relation between the wholes that share it** — load-bearing enough to carry sameness-attraction and the forceless recall connection.

That is more than the Law licenses on its own, and it should be given its own name so that a referee attacks the right thing. Two ordinary realists about universals can hold that two green things share one green without holding that the sharing constitutes a relation with consequences. The framework holds the stronger thing, and it needs the stronger thing: the whole account of recall, of sameness-attraction, and of how two systems can be alike without anything passing between them, rests on it.

Vocabulary discipline, ruled by the author (21 August). *True, not real. "Relations are true or false, and extrinsic or intrinsic, and hold some sameness and some difference, but reality I take to equate with the common view of material reality, physical things. So relations are matters of fact, whereas real things have properties that include them in the classical present Now."* The same discipline that keeps colour true-but-not-real.

2.4 Why the strictness matters downstream

The strictness is not decorative. It is what makes the recall connection forceless and distance-independent: sameness by Leibniz is not a force that attenuates but an identity that holds. Weaken it to similarity and the framework loses the mechanism, and has to reintroduce a carrier it has spent Tier 1 refusing.

This is why §3.4's *exact but partial* reading matters so much, and why the alternative — that the match is a matter of degree — could not be accepted.

3. What the icon is indiscernible from

This section corrects the earlier statement of the core claim. D9 as written names the wrong relatum, and the author's ruling of 21 August (F55) fixes it.

3.1 Not the object as a real whole

The pear on the table is a real object: it has mass, solidity, and whatever accelerations reality confers. The icon has none of these. Whatever the icon is indiscernible from, it is not the pear considered as *water-of-earth* — the whole real thing with its boosted properties.

3.2 Not the object's own inner state

Nor is it the pear's own internality — its Q3 state, the inner world of a pear that is just the relations among its parts. The author's ruling: the observer *"cannot help but form an icon, because it is not the actual object being observed."* The observer never has access to the pear's inside.

3.3 The observed's extrinsic air-from-earth

What the icon matches is the icon the object *presents*: **the extrinsic air-from-earth** — the relations among the object's parts as seen from outside, with the outward address intact. The author's words: *"the observer's icon is a reflection of the extrinsic whole of the observed... identical with the icon created as extrinsic air from earth, to the extent that it is accurate."*

Two consequences.

D9's core-claim sentence should be restated. It currently says the icon is indiscernible from *its worldly object*. It should say the icon is indiscernible from **the extrinsic whole the object presents**. This is a queued edit awaiting ratification, not a change made.

And it explains why a brain can be identical to a table in any respect at all — because both hold a void of the same shape. The author's formulation from *On the nature of Icons*: *"each have their own internality, and it is the matching of internality that makes the icon possess some identity with some other thing."*

One care needed here, because two formulations a week apart look different. The matching-of-internality formulation and the air-from-earth ruling are consistent only if the internality being matched is the one specified by the object's *extrinsic* boundary, not the object's own Q3 state. F55 is the later and governing statement, and the paper should use its wording. The two are reconciled by §4.1: the hollow is specified by its boundary, so matching the hollow and matching the boundary points are the same claim from two ends.

3.4 Exact but partial: the sampling account

The match cannot be total, and the reason is physical: *"the sampling rate of photons."* An icon *"is almost never, if ever, wholly identical with the observed in any regard except potentially as the objective fact."*

Asked whether this weakens Leibniz, the author ruled: **exact but partial**. Each sampled point is an exact identity; the shortfall is in **coverage**, not fidelity. Approximate in extent, exact in every particular.

Three things follow, and they are among the most useful results in the corpus.

The star's points are the sampled points. The geometry stops being a stipulation about where content happens to concentrate and becomes a consequence of how the content arrived. The icon is a star because perception is sampling.

Error is a closure failure, not a fidelity failure. Correct samples wrongly bound into a whole. This is what F7 already said, what the mirage case requires, and — as the literature check of `testability_v001.md` §4.1 found — what the memory-conjunction-

error work independently reports: *“features of stored stimuli maintain some independence in memory and can be incorrectly combined.”*

And it generalises into a principle. The author, 22 August: *“It all comes back to Leibniz applying to the partial view of things as much as to the wholes.”* Exact-but-partial is not a repair to the sampling problem. It is how the Law works on anything less than a whole, which is everything.

4. The geometry: star, hollow, and boundary

4.1 Star and hologram are one structure read from two ends

The apparent incompatibility — the star centrifugal with its content at outward points, the hologram encoding a bulk on a boundary — dissolves once a starpoint is given two ends.

From *On the nature of Icons*: the starpoints are how *“an icon reaches out along many different dimensions to connect with other facts, each starpoint is a fact in itself about the icon, and it gains its meaning from the extrinsic connection that it has.”* Within the icon there is *“only empty space unless it happens to be filled by another icon.”*

So: the **outer** end of each point reaches out into other facts and is where meaning lives; the **inner** end terminates on the icon’s boundary; the void between is specified by exactly those terminations.

The correction the author made, and it matters. An earlier Claude formulation had the starpoints’ inner ends *depositing* information on the boundary. That is wrong. The Q2→Q3 rotation is a turn inwards that subtracts Q1 wholeness, leaving the subjective perspective on the parts, so the space discovered is *necessarily empty until populated by some necessity* — and **the information that defines the space is still outside, and reaches only as far as the external edges of the space.** The author’s version is the derived one; the earlier was constructed.

The Susskind reference stays structural. Boundary encodes bulk, and nothing more. Importing the physics would repeat the mistake Tier 1 corrected when it withdrew the inverse-square-law claim about individuals: a fact about flux in three dimensions, asserted of something else, reads as numerology, and a referee will spring it.

4.2 Concentration of measure: the star as a theorem

The mathematics here was checked on 25 August 2026 and is stated in the form that survives scrutiny.

icon_dynamics_v002.md §3.3 says a high-dimensional star is *“a high-surface low-volume shape that has its substance close to its boundary.”* That is not merely an image. It is a well-known result in high-dimensional geometry, and it should be cited as one.

The safe form is scale-free. Whatever the radius, almost all the volume of a high-dimensional body lies in a thin shell at its boundary. For the outer 5% shell of a unit n -ball: 14% of the volume at $n = 3$; 40% at $n = 10$; 92% at $n = 50$; **99.4% at $n = 100$.** As dimension rises, interior volume becomes negligible and substance is at the surface.

A vivid cousin, worth one sentence: in high dimensions almost all of a cube's volume is in its **corners**. The inscribed ball is 52% of the cube at $n = 3$, 0.25% at $n = 10$, and about 10^{-70} of it at $n = 100$. The substance really is out at the points.

The unsafe form, which the paper must avoid. The volume of the *unit* n -ball rises to a maximum at $n = 5$ (2, 3.14, 4.19, 4.93, **5.26**, 5.17, 4.72 ...) and then collapses toward zero. This is true, and it is tempting, and it is *unit-dependent* — it holds the radius at 1, which is a convention. Stated flat as *adding dimensions decreases volume*, it is attackable. State the shell result instead; it says what the framework needs and cannot be dismantled on a technicality.

4.3 The hollow is room

This section is Claude's, developed from the author's own observation of 25 August that nesting reduces the emptiness. It is a proposal awaiting his ruling.

If substance concentrates at the boundary as dimension rises, then the hollow is what remains where dimensions are *absent* — and absence of a dimension is exactly what the framework means by a feature being left undefined.

So the icon's emptiness has a reading beyond the geometric: **the hollow is the icon's under-determination, measured**. It is room. And room, on the proposal in `conceivability_v001.md` §2.2, is the hypodox — the modality in which a thing is conceivable rather than either nothing or fully actual.

Three consequences, if it holds.

Abstraction is emptiness. A general concept is a sparse star with a large hollow: few dimensions determined, much left open, and so applicable to many things. A particular is a dense star with almost none: *this* pear, with *this* green, in *this* position. The abstract/concrete distinction falls out of the geometry rather than being stipulated, and it explains why abstraction is *useful* — the room is what lets one icon serve many occasions.

Full population is full determination. An icon populated to the limit would have no room left, and by the second constraint of the conceivability criterion it would have passed out of conceivability into usability. That is the integer two: *fully defined* — *not merely conceivable, it is usable*.

And the socket is the same object under another description. A socket is a delineated void — a hollow with a shape and no content. On this reading the socket is not a special item the mind manufactures for recall; it is what any icon has in the region where its dimensions are undetermined, made definite at the edges by the cue. Which is why the same machinery serves perception, recall, deduction and imagination: they differ in what shapes the hollow, not in what a hollow is.

4.4 Nesting is not stuffing

Also Claude's, and the answer to an objection the previous section invites.

If a rich icon has almost no interior, where does a nested icon go?

Not into the middle. A nested icon does not sit in the hollow like a marble in a box; it brings its own points, and those points become part of the containing icon's boundary

structure. Nesting reduces emptiness **dimensionally, not by occupancy** — adding dimensions is *how* the hollow shrinks.

So more detail at the surface and less empty space are not two facts that happen to travel together. They are one fact, and F64's word for what the composite does to its constituents is the right one: **layered**, not filled.

This matters for the comparison with IIT ([positioning_iit_v001.md](#) §2.3), where the nesting claim is the sharpest point of contradiction between the two theories. It is worth being precise about what Iconism is claiming when it says icons nest: not that consciousness contains smaller pockets of consciousness sitting inside it, but that the boundary structure of a mind's icon is elaborated by the icons composing it.

5. What the brain adds

5.1 Why not a mirror

If every object has an internality and an air-from-earth icon, and if a mirror samples the extrinsic whole at a far higher rate than an eye, why is a mirror not aware?

The author's answer (F57): *"the table cannot be like anything other than a table, the mirror can only present a detailed image like a photograph, a book can only tell the story it contains. A brain on the other hand has billions of connections, and so can model the perceived image, in real time, and link to it a potentially enormous quantity of other facts, populating the hollow hologram... which is all a fancy way of saying the brain can create an empty icon and then fill it with recalled icons from memory."*

5.2 The socket is the differentiator

The sharp form. The differentiator is not fidelity, not storage, not connection-count. It is **the capacity to form a delineated void and fill it**. A mirror cannot create a void; it can only hold a correspondence. A book holds shapes that will occasion icons in a reader (F46) — meaning held in a medium is held inanimately, and the score is not the performance *even if it includes all the dots*.

This is Q3's own definition with a mechanism in it. Tier 1 gives Q3 as *"necessarily empty except where populated by other views from the inside."* The empty icon filled with recalled icons is that sentence with a mechanism.

And it says why dimensional reach is the criterion, which is the framework's answer to panpsychism in one move. Every object has an inside; the table's does not overlap its legs'; all of them are hollow. What the brain adds is that its hollow has the dimensional reach to be *like* what it perceives, recalls, imagines or deduces. The author's own formulation of the target, from August 2026: *a new way of looking that shows them why a table cannot be panpsychic, but a brain can be, and for exactly the same reason.*

5.3 The same answer, applied twice, with no verdict implied

Iconism's answer to *why not a mirror* and its answer to *why not a machine* are one answer applied twice: can the system form a delineated void and fill it?

No verdict on the second is implied, and none should be smuggled in here. The criterion is F49's — the link must be **virtual** (X is like Y), not **material** (X stands in for Y) — and the author's own framing of the machine case remains the Locked-In analogue at the substrate level: *"which might be sad if it means that you are conscious but cannot actively affect the output you create with any input from that consciousness."* O12 stands and is the author's to rule. `ai_and_sentience.md` develops the case.

6. Reference, and why this is not representationalism

6.1 Reference runs through icons

Words are pointers (F38). A word detached from an icon is salad. Reference is not a relation between a symbol and a thing; it is a relation between a symbol and an icon, which in turn holds the anchored relation to the thing. Two links, not one, and the second is where all the work happens.

This is why the framework's account of conceivability can say that a stipulation for which no icon can be formed is *a form of words with an illusion of meaning* (`conceivability_v001.md` §3), and why the Chinese Room reformulation bites: *"a physical aping of the mental process of mentally aping a physical process"* — one level of mimicry too many.

6.2 How this differs from representational accounts

Standard representationalism has a mental state *stand for* a worldly state, with the standing-for relation grounded in causal covariance, teleofunction, or asymmetric dependence. On every such account the vehicle need not be *like* what it represents at all — a token is a token, and its arbitrariness is usually counted a virtue.

Iconism holds the opposite on the shared dimensions. The icon is not a token standing in for the pear; it is indiscernible from the pear's extrinsic whole at the sampled points, and by Leibniz's Law the features there are *the same features*. That is F49's distinction: the material link is blind, the virtual link holds because X is like Y.

The framework should not overclaim here. It is not offering a rival theory of mental content to sit beside Dretske, Millikan and Fodor. It is making a narrower and stranger claim: that on some dimensions there is no representation relation to theorise, because there are not two features to relate.

6.3 The two objections that have to be answered

These are the objections that have sunk likeness accounts before, and a paper that does not meet them will not be read past the point where they occur to a referee.

Berkeley: an idea can be like nothing but an idea. The objection assumes that ideas and objects are different kinds of stuff, so that any likeness between them is a category mistake. The framework denies the antecedent. There are not two substances; there is one substrate and layers of fact about it (`icon_dynamics_v002.md` §8). The shared feature is neither mental nor physical — it is a *fact*, a true relation, one green instantiated twice. Berkeley's objection is a good objection to a dualist likeness theory, and Iconism is not one.

Goodman: resemblance is cheap. Similarity is reflexive, symmetric, and holds between any two things in some respect or other; so it explains nothing about why this represents that. This is the harder objection and it needs three answers, which together are enough.

- *It is not similarity.* The framework's relation is indiscernibility at a point, which is identity of a feature, not resemblance of a whole. Goodman's strictures are aimed at a relation that admits of degree; this one does not.
- *It is not cheap, because the points are the sampled points.* What makes an icon the icon of *this* pear rather than of pears in general is its **coverage** — which points are matched and how many. Selectivity is doing the work that similarity could not (§3.4).
- *And the direction is not supplied by the likeness.* See below.

6.4 Direction: what makes the icon of the pear rather than the pear of the icon

Indiscernibility is symmetric. Aboutness is not. So something other than the likeness must fix which of the two is the icon, and the framework has two answers that do not compete.

The causal-historical answer. The photons went one way. The icon is built from samples of the object's extrinsic whole; the object is not built from samples of the icon. Direction comes from the sampling (§3.4), not from the relation.

The structural answer, and it is the better one. The icon is the item with the **hollow that was shaped to be filled**. The pear has no socket. What makes something an icon of is that its void was specified — by cue, by input, by the facts active in the region — and then occupied. Aboutness is a fact about which side did the specifying.

Claude's note. This is the framework's most direct answer to the objection that likeness cannot ground intentionality, and as far as I can find it is not stated anywhere in the corpus in this form. If it holds, it should be stated in the paper in a sentence, because the objection is otherwise the first thing a philosopher of mind will reach for.

7. Error, fiction, and the limits of anchoring

7.1 Mis-occupancy

The void is **permissive, not selective**: *"anything that fits in the intrinsic void might float in there, which I think is why we misremember so often"* (F59). Error is the wrong icon floating into a void that happens to fit it — a closure failure, with the parts intact.

That is also why the framework needs **no negation** (F60). Icons cannot be negated: to negate one, each relation would have to be cancelled, and the similarity cannot be cancelled because it is the very dimension of the relation, while a degree of difference cannot be negated either. Wrongness was never negation; it is mis-occupancy, and what the framework needs is a fit criterion, which it has. `conceivability_v001.md` §6 develops this and shows it is the same structure as inconceivability at another scale.

And the void is not a negative. *"The true void is not a negation of being but a reflection of it, there is no negative to being, a thing either is, or it is not. A negative is a thing turned*

on its head, not the subtraction of the thing” (F58). An intrinsic void is surrounded by whatever defines it as intrinsic, which is why the delineated void is *defined* and the undelineated is merely not-what-exists.

7.2 The hallucination taxonomy

Three failures, at three different places, and distinguishing them is a strength of the account rather than a complication.

- **Brain-guessing** — layer 3. The mind settles on an apparent fit; the socket is filled by a near-match brightened too soon (F7). Misidentification and misremembering live here.
- **Interference in the mechanics** — layers 1–2. A drug such as LSD disturbs the parts-supply, and the icons that form are icons of nothing that was sampled.
- **Interference in the environment** — outside the system entirely. A heat haze produces a mirage, and **this is not an error of mind at all**: the icon is a correct icon of what the light actually did, and the mistake is downstream, in what is inferred about the source.

The mirage case is worth stating explicitly in the paper because it looks like a counterexample and is in fact a confirmation: the anchoring held: the icon matched what reached the observer.

7.3 Fiction, imagination, and the departure from Hume

Along the extension away from the primal case the anchor weakens. **Fiction is the limit: a partially-untrue fact** — an icon with created components and a worldly anchor weakened to nothing (F22). It is not a failure of truth but a different kind of fact, and it is what propaganda exploits: a partially-untrue fact offered as though fully anchored (F69).

Hume held that imagination can only recombine what sense supplied. The author departs from this: *“the thing imagined may be entirely original... but even then the icon has to be constructed.”*

Claude’s proposal, flagged in the 21 August dialogue and unratified. The departure may not need a special faculty. The author’s own analogy has an icon’s interior content forming *“in exactly the same fashion that ratios form as the roots of natural numbers on the number line”*, with the rider **not all roots are regular**. Read across: the shared, anchoring dimensions are the *commensurable* interior content, expressible in terms of the boundary data; the created dimensions are the *incommensurable* — forced by the boundary, not expressible as a ratio of it. Magenta is the irrational root.

If that holds it does two things. It gives the shared/created split a mathematical form instead of a definitional one, using Tier 1’s own number line. And it explains the Hume departure without a faculty: irrational roots are *generated*, not copied, so an icon can carry interior content that was on no prior boundary.

8. Multiple realisation

8.1 Putnam and Fodor

The multiple-realisation argument holds that mental kinds cannot be identified with physical kinds because the same mental state can be realised in indefinitely many physical ways.

Iconism agrees and is not embarrassed, because **the icon is individuated by its layer-3 relation** — indiscernibility with an extrinsic whole at the sampled points — and not by the layer-2 substrate that supports it. The same icon can be produced by different firing patterns; a given firing pattern in a different brain or a different context would not produce the same icon.

The framework's position is therefore **stronger than functionalism and more permissive than any identity theory**. Stronger, because it says *what* the common feature is — a whole indiscernible from a worldly whole on the dimensions shared — rather than merely that something common exists. More permissive, because the substrate is not what individuates.

8.2 Kim

Kim's causal exclusion challenge is answered in `icon_dynamics_v002.md` §5 by the wave-function-contribution mechanism, and is not re-argued here. What belongs here is only the individuation point: the reason layer-3 facts are not redundant with layer-2 facts is that they are facts *about relations among* layer-2 facts, and those relations are what the icon is.

8.3 Sorites, not Theseus — and the invariance property

Tier 1 supplies the model and Iconism may rely on it: facticity is preserved even where the truth-makers are recycled, and the model is **Sorites rather than Theseus**. Theseus asks whether it is the same ship after replacement, and the question only arises if identity is carried by the parts. On this account it never was. *The parts halve and halve again without end, and the heart does not move.*

This is the same property Tier 1 specified for the intrinsic void of individuality — *invariant under the replacement of everything that determines it* — and the convergence is worth stating: what is multiple realisation seen from the Iconism side is Sorites-not-Theseus seen from the Tier 1 side, and it is one property under two descriptions. An icon does not cease to be the icon it is because its substrate has turned over.

9. The Cartesian theatre, reframed

Dennett's denial is granted in full: there is no place in the brain where icons are presented to an audience. No theatre, no homunculus, no seat.

What the framework refuses is the inference from *no place* to *no fact*. Icons are **abstract in the strict sense** — they have no specific position in the brain because their correlates are dispersed throughout it — and what does the staging is not a location but a **configuration's focus**: the icon most strongly and meaningfully connected to the rest

is the brightest, is where the centre sits, and is what gets staged. Centring just *is* brightness by meaningful connection (F17).

The theatre is not a room. It is real on the same as-if basis a centre of gravity is real. `centre_of_awareness.md` develops the analysis; what belongs here is only that the abstractness of icons is what makes the Dennett concession costless.

10. What this document does not do

- **It does not develop the centre of awareness.** §9 states the reframing; `centre_of_awareness.md` gives the account.
 - **It does not settle the AI question.** §5.3 states the criterion and stops.
 - **It does not argue the wave-function-constraint mechanism.** That is `icon_dynamics_v002.md` §5.
 - **It does not develop the three-way consciousness/awareness/spirit distinction,** which is `consciousness_awareness_spirit.md`'s.
 - **It does not offer a theory of mental content** to sit alongside Dretske, Millikan and Fodor. It makes a narrower claim (§6.2) and should be held to it.
 - **It does not take a position on the eternal-fact ontology** that the persistence of icons after demerger would require. Spirit's.
-

11. Status and open questions

Firm. The two relations (§1.2). Leibniz on features (§2). The air-from-earth relatum and exact-but-partial (§3, F55, F56). Star and hologram as one structure, with the author's correction (§4.1). The socket as differentiator (§5). Reference through icons; meaning in a medium inanimate (§6.1). Mis-occupancy and the permissive void (§7.1). Multiple realisation (§8).

Queued edits to `icon_dynamics_v002.md`, awaiting ratification.

- **D9's core-claim sentence** names the wrong relatum and should read *the extrinsic whole the object presents* (§3.3).
- **§3.4's location of indiscernibility.** The document puts it *at the points*; the matching-of-internality formulation puts it in the shape of the void. §4.1 reconciles them — the hollow is specified by its boundary — but the wording should say so rather than leave a reader to notice.
- **§3.3 gains the concentration-of-measure result** (§4.2), in the scale-free form only.

Proposals awaiting the author's ruling.

- **§4.3 — the hollow is room.** Emptiness as under-determination, and the abstract/concrete distinction falling out of it. Claude's, developed from the author's own 25 August observation.
- **§4.4 — nesting is not stuffing.** Claude's.

- **§6.4 — direction supplied by the socket.** Claude's, and the most consequential of the three, because it answers the standing objection that likeness cannot ground intentionality.
- **§7.3 — magenta as the irrational root.** Claude's, flagged unratiſed ſince 21 Auguſt.

Open.

- **O11** — do wholes of meaning obtain in arrangements of atoms as well as in brain-activity? This document assumes nothing either way; §5.2's dimensional-reach criterion is compatible with both readings, which is either a virtue or an evasion depending on how O11 is ſettled.
- **The boundary question.** If the whole information content ſits at the boundary, what fixes the boundary — what makes *theſe* ſtarpoints one icon rather than two? F19 (ſiloing) and F20 (forced emergence) answer it, but the holographic reading raises the ſtakes, and this is the p-vampire's own question asked from inside the geometry. It needs ſtating at length ſomewhere, and this document is the natural home.
- **Whether §6.2's modesty is the right call.** The framework may have more to ſay againſt Dretske and Millikan than "we are making a narrower claim." Not attempted here.

12. Candidate aphorisms

- *"I uſe icon in the religious ſenſe, i.e. that the image is ſo like the reality that the two have the ſame factual properties."* (Auguſt 2026.) The core claim for a reader with none of the apparatus, and the answer to why the word is *icon*.
- *"Not two ſimilar greens but one green."* Leibniz on features, compressed.
- *"It all comes back to Leibniz applying to the partial view of things as much as to the wholes."* (22 Auguſt 2026.) Exact-but-partial as a principle.
- *"The table cannot be like anything other than a table."* (21 Auguſt 2026.) The whole of §5 in eight words.
- *"The ſcore is not the performance, even if it includes all the dots."* (F46.) Meaning in a medium is inanimate.
- *"The true void is not a negation of being but a reflection of it."* (F58.)
- *"Anything that fits in the intrinsic void might float in there."* (F59.) Error as mis-occupancy.
- *"The parts halve and halve again without end, and the heart does not move."* (Tier 1.) Sorites, not Theſeus — borrowed, and worth borrowing.

Companion documents: *icon_dynamics_v002.md* (§3 the geometry, §5 downward cauſation, §8 Princess Elisabeth), *conceivability_v001.md* (§2.2 the determination ſcale, §6 wrongneſs without negation), *positioning_iit_v001.md* (§2.2 aboutneſs, §2.3 neſting), *testability_v001.md* (§4.1 the memory-conjunction-error convergence), *Iconiſm Interface v002* (§1 Q3, §4 Sorites-not-Theſeus, §5 the commitments). Next in the chain: *conſciouſneſs_awareneſs_spirit.md*, then *centre_of_awareneſs.md*, then *ai_and_sentience.md*.