

The Axiom of Hope: Indexical Realism and the Structure-Dependence Problem

Nathan G. Whittingham

Preprint · v1.0 · 27 January 2026

Stated manuscript version date: 27 January 2026. Web and PDF presentation prepared 5 September 2026.

Context for this version

This preprint preserves the argument presented as v1.0, dated 27 January 2026. Subsequent research has refined and narrowed aspects of the surrounding framework, while this manuscript’s argument is preserved. The Research page describes the current direction of the research.

Abstract

Luciano Floridi’s informational structural realism (ISR) holds that, as far as we can tell, reality is informationally structured and that invariants across Levels of Abstraction (LoAs) give us non-arbitrary grip on it. I grant invariants as our best available evidence, but challenge the step from cross-LoA stability to knowledge of mind-independent reality. Under present epistemic conditions, every usable truth test comes from within a framework; none stands outside all LoAs. Hence the Structure-Dependence Problem (SDP): without a correctness norm, outputs are not knowledge; with one, methods remain LoA-bound; and no LoA-external truth test is available. ISR’s realist conclusion therefore rests on an axiom of hope—sound as a research policy, not yet a result. I distinguish a policy reading (Policy-ISR) from a result reading (Result-ISR), defend the former, and reject the latter. The constructive upshot is Indexical Structural Realism (IxSR): treat invariants as warrant within an admissible family of LoAs, and turn framework-dependence into method. I operationalize IxSR for AI and digital humanities—where pipelines are explicitly LoA-bound—with practical guidelines (multi-LoA design, interoperability, robustness sweeps, “index” lines for claims, and independence/translation-debt reporting). The contribution is both diagnostic and constructive: a principled boundary for realist claims today and a design discipline for building more reliable knowledge systems.

Keywords: Informational structural realism (ISR); Levels of abstraction (LoA); Structure-Dependence Problem (SDP); Indexical structural realism (IxSR); Robustness analysis; Digital humanities and AI

1. Introduction — Framing the Debate

Luciano Floridi’s informational structural realism (ISR) holds that reality, as far as we can tell, is informational in structure, and that invariants across Levels of Abstraction (LoAs) give us non-arbitrary traction on it. I grant that this is our strongest available evidence. My claim targets one step only: the inference from cross-LoA stability to knowledge of mind-independent Reality. Under current conditions, we lack any truth test that stands outside all frameworks, so treating invariants as correspondence ultimately depends on an axiom of hope—a sensible research policy, not yet a result.

The **Structure-Dependence Problem (SDP)** follows: without a correctness norm there is no knowledge; with one, the method is LoA-bound; and no LoA-external truth test is available. The same limit governs digital humanities (DH) and Artificial Intelligence (AI), whose pipelines explicitly fix observables, inferences, and norms. This paper formalizes the dilemma for ISR and develops **Indexical Structural Realism (IxSR)**—a design method for robust, transparent, interoperable knowledge systems within acknowledged bounds.

This paper makes two moves. First, it shows that ISR’s route from invariants to realism depends on a truth-test we do not have; without it, ISR slides back toward internal realism. Second, it offers a constructive path: **IxSR**. On this view, invariants are valuable **within** bounds. The task is to build robust, transparent, and interoperable knowledge systems by **designing for multiple frameworks, publishing the bridges** between them, **sweeping robustness, showing the frame, and scoping the claim**. That stance fits the realities of DH and AI practice and avoids metaphysical overreach.

The sections ahead situate the argument in the literature, spell out why ISR’s evidence is not yet enough, and translate the lesson into concrete DH/AI guidelines.

2. Relation to Existing Literature

Floridi’s ISR sits at the intersection of structural realism in the philosophy of science and the broader realism/anti-realism debate in epistemology. It combines an **ontological thesis** — that the ultimate nature of reality, *as far as we can tell*, is informationally structured — with an **epistemological thesis** — that all access to reality is mediated by LoAs. In *The Philosophy of Information* (2011) and *The Logic of Information* (2019), Floridi develops LoAs as modelling interfaces that fix what counts as an observable, determine permissible inferences, and set the norms for correctness. His account of **semantic information** as “well-formed, meaningful, and truthful data” links informational content directly to truth conditions. ISR’s promise is that, while we cannot take a “God’s eye” view, we can still capture reality’s structure through the invariants that survive across LoAs.

This combination has drawn interest and critique, though none directly replicate the Structure-Dependence Problem presented here.

2.1 Structural Realism and Its Limits

Floridi explicitly builds on **John Worrall's** (1989) **epistemic structural realism** (ESR), which holds that what survives radical theory change is not the objects posited by scientific theories, but their structural relations. Worrall's view is a compromise: it preserves the realist intuition that science tracks something real, while avoiding the pessimistic meta-induction that undermines belief in specific theoretical entities. Worrall, however, does not address how such structures are accessed or fixed, leaving him open to questions about framework dependence. Floridi's LoA formalism is meant to supply that machinery.

Yet, as I argue here, LoAs themselves create an epistemic bind: either a method for accessing reality lacks a correctness norm (and so is not knowledge), or it has one, in which case it remains LoA-bound. This conditional dilemma is not present in Worrall's own work or in the subsequent ESR literature. Worrall treats structure as the stable content of science without interrogating whether the identification of structure is itself framework-relative.

2.2 Realism, Anti-Realism, and Conceptual Schemes

Hilary Putnam's (1981) **internal realism** rejects the "God's eye" view and treats truth as idealised rational acceptability within a conceptual scheme. On this view, all knowledge claims are framework-dependent; there is no standpoint from which to compare our conceptual schemes to reality "as it is in itself." **Bas van Fraassen's** (1980) **constructive empiricism** is similarly cautious, restricting belief to the claim that science's observable consequences are true, without committing to the truth of unobservable entities or structures.

Floridi departs from both: ISR is meant to be realist "as far as we can tell." By appealing to invariants across LoAs, he hopes to secure knowledge of reality's structure in a way that internal realism and constructive empiricism would reject. My argument shows that without escaping framework-bound verification, ISR collapses into something much closer to internal realism — a position Floridi explicitly resists.

2.3 Mediation and Epistemic Opacity

In the philosophy of technology and computational science, **Paul Humphreys** (2004) has explored **epistemic opacity** — the idea that complex computational processes can generate knowledge claims whose inner workings are inaccessible to human cognitive agents. Humphreys stresses that all computational models operate within representational constraints, determined by modelling choices, approximations, and algorithmic design. While Humphreys' work aligns with my claim that knowledge is mediated by frameworks, he does not target Floridi's ISR or articulate the strict either/or dilemma that underpins the Structure-Dependence Problem.

2.4 Digital Epistemology and the Humanities

In the DH, **Johanna Drucker** (2011) has argued that data should be understood as *capta* — constructed through interpretive frameworks rather than simply “found.” **Geoffrey Rockwell** and **Stéfan Sinclair** (2016) similarly emphasise the interpretive loops embedded in text-analysis tools, while **Ted Underwood** (2019) shows how large-scale machine learning methods in literary studies depend on decisions about features, corpora, and classification schemes. These accounts resonate with the idea that digital knowledge production is LoA-bound: every computational pipeline instantiates its own admissible observables, inference rules, and correctness norms.

What none of these works attempt, however, is to connect the epistemic situation of DH methods to the ontological and epistemological commitments of ISR, or to show how such commitments are vulnerable to a dilemma that arises even if one grants Floridi’s ontology.

2.5 Where This Argument Extends the Literature

This paper accepts Floridi’s ontology for the sake of argument and shows that verification remains trapped within LoAs—a move largely absent from existing critiques. It frames the issue as a strict dilemma: without a correctness norm there is no knowledge, while with one the method is LoA-bound. And it recasts Floridi’s invariants as yielding only indexical progress—stability relative to admissible frameworks—grounded in DH/AI case studies rather than asserted in the abstract.

By connecting the LoA apparatus, the invariants strategy, and the epistemology of digital systems, this argument shows that the same limits apply to AI-supported knowledge production in the humanities as to ISR’s own claim about knowing reality’s structure. In doing so, it extends the philosophy of information into a domain where these epistemic questions are not just abstract puzzles but practical constraints on research design and interpretation. My target is Floridi’s informational structural realism—where semantic information is truth-apt and LoAs mediate access—not ontic structural realism in general.

A sympathetic reading may say that my “Indexical Structural Realism” merely restates ISR in modest terms. The difference is not cosmetic. ISR, as often presented, treats cross-LoA invariants as evidence that we are latching onto reality’s structure; IxSR restricts the epistemic claim to **indexical warrant**: stability relative to an admissible family of LoAs and published translations. IxSR adds enforceable method: independence criteria, translation maps, a robustness sweep, and a claim-scope rule. If ISR is read as a **research policy**, we largely agree on practice; if it is read as a **metaphysical result** now justified by invariants, we part ways. The dispute is about the status of the inference, not about the value of invariants.

Robustness analysis in philosophy of science already recommends seeking results that persist across independent models and measurements (Levins 1966; Wimsatt 1981; Weisberg 2006/2013; Steel 2008). IxSR inherits that spirit but adds Floridi’s LoA formalism and enforceable practice: explicit translation maps, independence audits across data/representation/metric, and claim-scope statements. A second, classic worry is Newman’s problem (Newman 1928; see Worrall 1989): if “structure” is too thin, any set with the right cardinality can realize it. IxSR avoids triviality by requiring role-preserving

translations¹ (preserving measurement or causal roles, not mere relabeling) plus independence criteria that rule out pseudo-diversity. The result is a constrained notion of invariant that does real evidential work without collapsing into set-theoretic vacuity.

3. The Present Epistemic Condition — Why Floridi’s Claim Isn’t Supported (Now)

Informational Structural Realism pairs a strong ontological thesis with a careful epistemology. Ontologically, it says reality’s ultimate nature is informational. Epistemologically, it holds that all access is mediated by Levels of Abstraction—structured interfaces that fix what can be observed, how observations relate, and which results count as correct. The question is whether this pairing justifies Floridi’s realism under our present epistemic condition. It does not.

By the present epistemic condition I mean two facts. Every usable truth test we have comes from within a framework; each test arrives with rules and instruments attached—measurement protocols, accuracy metrics, statistical thresholds, calibration routines. And we have no truth test that stands outside every framework. There is no way, under current conditions, to verify claims about mind-independent reality without relying on some framework’s standards. ISR asks invariants to carry more weight than that setting allows.

A framework, as I use the term, is simply a package of commitments about what can be seen, how to infer from it, and how to judge correctness. Floridi’s LoAs are a special case: frameworks fixed by a modelling interface. A correctness norm is any standard that marks right from wrong—truth conditions, accuracy metrics, calibration targets. An invariant is a feature that persists when one switches among an admitted set of frameworks, preserved by translations that keep measurement or causal roles rather than by mere relabelling². An outcome is indexical when it is true or stable only relative to the stated set of frameworks. By “admissible,” I mean frameworks that are (i) independent in data, representation, or metric; (ii) linked by explicit, published translations; and (iii) candid about the norms they use to score success³.

Floridi’s view rests on four ideas. Reality is informational in nature. All knowledge claims arise within LoAs. Semantic information is truth-apt—content as “truthful data,” not mere syntax. And invariants across LoAs are supposed to track the world’s real structure. That last step needs a truth test not fixed by any LoA. We do not have one.

On truth, Floridi has two options. If he deflates truth to “what our best framework would endorse,” ISR slides into internal realism. My objection loses bite, but so does the realism that gives ISR its appeal. If he keeps truth as correspondence to how things are beyond our frameworks, then the verification step cannot be fixed by any framework—and no such step is available now. Call this the truth fork. Either way, the strong realist claim overreaches.

There is no escaping the frame. Calibration needs a standard; choosing one already places us inside a framework. Measurement is rule-governed; to call an outcome correct is to follow agreed procedures. Translation between frameworks requires criteria of equivalence, and those criteria are themselves

framework-bound. Even expert consensus is not an exit; it is agreement under shared standards.

LoAs are valuable because they make these commitments explicit. They also expose the verification gap. If every claim is generated and judged inside some LoA, then any realist truth-check would have to stand outside them. ISR denies any “view from nowhere,” yet still asks invariants to close that gap. They cannot do so without borrowing a test from somewhere, and all the “somewheres” we have are LoAs.

Consider a breakthrough method that claims to know reality “as it is.” By what rule could we tell it is correct? Without a rule, it is not knowledge—there is no way to distinguish success from failure. With a rule, the method stands within a framework that supplies it. Either way, the result does not outrun the framework.

A strong realism claim would be warranted only if a correctness norm exists, that norm is not supplied by any framework, and it can be used to check outputs. At present, the second condition fails; all usable norms come from within frameworks. The conclusion follows: claims that outrun frameworks overreach.

Predictable replies do not change this. Reliabilism says a process can yield knowledge if it is reliable even if we cannot state the rule. But reliability must still be assessed, and assessment uses norms. Interventionism says successful manipulation shows reality. Yet success is scored by aims, protocols, and thresholds we accept. Deflationism makes truth whatever our best framework would endorse—a coherent stance, but not the robust realism ISR advertises.

Invariants can license use. If a pattern persists when we change tokenisation, swap models, or remap an ontology, the result is sturdier and worth trusting in practice. But invariants cannot, by themselves, license claims about reality “as it is in itself.” Endurance is always checked by translations we approve, and those approvals come with rules. The result is indexical: robust relative to the chosen frameworks and translations.

We can raise or lower our confidence by increasing the independence of the frameworks we test—drawing on different data sources, different modelling families, different success criteria, different preprocessing, different annotation guidelines, and different mapping rules. Breadth with independence is worth more than breadth alone; ten near-identical setups do not equal two genuinely independent ones⁴.

The point is concrete. A thermometer reads “right” because it was calibrated to a standard; change the standard and the reading changes. A theme that appears in both topic models and embeddings after harmonising vocabularies is stronger than one model alone, yet still tied to the preprocessing, corpora, and mappings used. Cross-collection answers in a museum database may stabilise after mapping to a common ontology, but they remain products of that ontology and its constraints.

The fair-minded policy is simple: treat cross-LoA invariants as provisional indicators of target structure, and adjust confidence by the breadth and independence of the frameworks tested. That is an honest, productive way to steer inquiry. It is not, by itself, a warrant for saying we know reality “as it is in itself.”

These same limits shape digital humanities and AI. Their pipelines are made of LoAs, their success metrics are correctness norms, and their “invariants” are patterns that survive across tools or models. The next section turns to practice: how structure-dependence shows up in DH and AI, and how to design for robustness without pretending to have escaped the frame.

4. Digital Humanities as LoA-Bound Systems

These abstract concerns about LoAs are not limited to philosophical speculation. They are embedded in the daily practice of research across many fields—especially in the digital humanities and AI—where design choices define what can be observed, how it can be processed, and what counts as a correct result. In these contexts, LoAs are not philosophical abstractions but operational realities that shape the epistemic boundaries within which all results are produced.

4.1. NLP in Literary Analysis

Consider topic modelling on a corpus of nineteenth-century novels. One LoA uses stemming and a standard stop-word list; another uses lemmatisation and a custom stop-word list developed to preserve function words common in Victorian prose. Even before modelling, the set of “observables” differs.

Apply Latent Dirichlet Allocation (LDA) to the first LoA and a neural embedding model to the second. A theme—say, industrialisation—might appear in both. In LDA it emerges as a probabilistic cluster of high-frequency terms; in the embeddings it appears as a dense region in a semantic vector space.

This convergence is robust within these LoAs, judged by topic coherence (LDA) and cosine similarity (embeddings), but the “same theme” judgment depends on those model-specific norms and preprocessing choices; change them and the convergence may disappear.

4.2. Computer Vision in Archives

Handwritten Text Recognition (HTR) systems transcribe medieval manuscripts using training datasets, image segmentation rules, and correctness norms. One LoA trains on transcriptions from one paleographer’s guidelines; another uses a different expert community with alternative interpretations for ambiguous letterforms and ligatures.

Run the same HTR architecture on both training sets. The model may produce identical readings for many words, giving the appearance of invariance.

The identical readings yield matching CER/WER (character-/word-error rate), yet “ground truth” differs by guideline, so agreement is correctness-relative, not manuscript-absolute.

4.3. Cultural Heritage Databases and Metadata Ontologies

Museum databases often catalogue artefacts using distinct metadata ontologies. One LoA uses a simple object-based schema; another maps that schema into the event-based CIDOC-CRM ontology. Mapping requires interpretive choices—whether a location change is an “event,” whether creation is a single act

or a chain of sub-events.

Run the same cross-collection query in both systems. Some results—say, all objects by a specific artist—will align.

The alignment holds under CIDOC-CRM and the mapped source schema, validated by SHACL (W3C Shapes Constraint Language) or competency questions, but that stability is tied to the chosen ontologies and mapping criteria.

In each case, cross-LoA invariants are valuable for coordination and reliability, yet their warrant remains indexical: stable only relative to the admissible frameworks and norms that produced them. This is the Structure-Dependence Problem in operational form: DH and AI systems can achieve robustness, but that robustness is always tethered to the design choices and evaluative standards of their LoAs.

These cases show the Structure-Dependence Problem in practice: even when results are stable across markedly different tools or models, that stability is anchored to the norms and translations fixed by their LoAs. The pattern is not unique to digital systems—it reflects a deeper philosophical inheritance. The next section situates this operational reality within a longer lineage, from Kant’s phenomenal/noumenal divide to Worrall’s structural realism, showing how Floridi’s LoA framework draws on these traditions while inheriting their limits.

5. Historical Context and Lineage in Philosophy of Science

The Structure-Dependence Problem visible in these DH and AI cases is not unique to digital systems. It reflects a deeper philosophical inheritance: long-standing attempts to explain how knowledge can be both successful and yet mediated by conceptual or structural constraints. From Kant to contemporary structural realists, the same tension recurs—how to account for stability in what we know without claiming a God’s-eye view. Floridi’s LoA framework draws on this lineage and inherits both its strengths and its limits.

5.1 Kant: Phenomenal and Noumenal

Kant preserved empirical knowledge by confining it to the phenomenal world —reality as structured by our cognitive faculties—and giving up claims to know the noumenal world “as it is in itself.” LoAs echo this in making all access framework-dependent. Where Floridi departs is in claiming that cross-LoA invariants can approach the noumenal structure, a step Kant would have rejected.

5.2 Putnam: Internal Realism

Putnam collapsed the noumenal entirely, holding that truth is idealised rational acceptability within a conceptual scheme. On this view, all knowledge is framework-internal. Floridi positions ISR against this, aiming to secure truth about reality’s structure via invariants. Without a non-LoA-bound verification, however, ISR risks sliding back into Putnam’s position.

5.3 Worrall: Epistemic Structural Realism

Worrall argued that science tracks structure because structural relations survive theory change. Floridi builds on this, replacing “theory” with “LoA” to capture structure in a more formal, general way. My Indexical SR accepts the value of invariants but makes their framework-relativity explicit, reintroducing the verification gap that Worrall leaves unresolved.

5.4 The Digital Turn

In digital humanities and AI, “frameworks” are not just conceptual but computational and infrastructural: tokenisation schemes, metadata ontologies, training datasets, evaluation metrics. Large-scale infrastructures like Europeana or the HathiTrust Research Center produce robust invariants across systems, but always within the bounds of their design choices. This is the same structure-dependence Floridi inherits: invariants are valuable, but their warrant remains indexical to the admissible frameworks in which they are found.

6. The Structure-Dependence Problem — The Conceptual Challenge

6.1 Frameworks and the Dilemma

Theorem (Structure-Dependence). For any method M: (i) without a correctness norm, M’s outputs are not knowledge; (ii) with a correctness norm, M is bound to the LoA that fixes that norm; (iii) no LoA-external truth test is available under current conditions. **Corollary.** Cross-LoA invariants warrant **indexical robustness**—stability relative to an admissible family of LoAs and their published translations—but do not, by themselves, license correspondence claims about reality “as it is in itself.”

Bridging Principle (B). If a property P is invariant across an admissible family of LoAs, then P corresponds to target structure.

Comment. B is not entailed by the definitions of LoA, correctness norm, or invariance. Either we adopt B inside a meta-LoA (so its warrant remains LoA-bound) or we adopt B as a research policy (the axiom of hope). Without B, invariants support use; with B, the realist leap is justified only by the very framework that scores the invariants.

As defined in §III, a framework fixes what counts as an observable, determines which inferences are permissible, and provides a norm for correctness. On Floridi’s own account, to possess semantic information is to meet a truth-linked norm. In other words, something counts as knowledge when it succeeds according to a recognised rule of correctness.

This is where the tension with ISR is sharpest. Floridi hopes to secure realism “as far as we can tell” by tracking invariants across multiple LoAs. But the inference from invariants to correspondence with ultimate structure requires a verification step that is not fixed by any LoA. Without that step, the most the invariants can deliver is what Indexical Structural Realism openly acknowledges: stability indexed to the set of admissible frameworks and the translations we have chosen.

Framed this way, the Structure-Dependence Problem is not an eccentric objection. It formalises the gap that ISR inherits from earlier attempts to reconcile realism with mediation: the Kantian limit on the noumenal, the internal realist's confinement of truth to a scheme, and the structural realist's reliance on preserved relations without independent access to the relata. ISR's LoA apparatus makes the commitments explicit, but it does not supply the missing non-LoA truth-test.

The rest of this section addresses possible escape routes—reliabilism, interventionism, and the “all-information” reply—and shows why none of them dissolve the dilemma.

6.2 Replies and Refinements

Reliabilism. A reliabilist might argue that a process can yield knowledge if it is reliable, even without access to the rule. But reliability must still be assessed, and that assessment depends on norms drawn from some framework. Calling a process “reliable” without identifying the standard simply hides the framework that makes the label possible.

Interventionism. Interventionists hold that if we can manipulate something successfully, it is real. Yet success is always scored by aims, protocols, and thresholds we accept—each defined within a framework. Push-back from the world strengthens warrant inside that frame; it does not supply a framework-free test.

The ‘All-Information’ Reply. Floridi might say that if reality and LoAs are both informational, the mediation problem is dissolved—like knows like. On this view, the framework is not an alien imposition but a way of tuning into a particular resolution of the underlying informational reality. Even so, the LoA's specific observables, granularities, and correctness norms remain a contingent, mediating structure. They fix how the informational world is accessed and scored, and those design choices are not dictated uniquely by the world's own structure. In an all-information world as in any other, we still need a way to distinguish genuine structure from artefacts of the modelling interface, and that discrimination is itself governed by the LoA. The shared ontological category does not remove the need for a framework; it only makes its homogeneity part of the background, leaving the correspondence between LoA and ultimate structure unverifiable from within.

6.3 Mini-Cases

Consider temperature invariance: the concept appears in both thermodynamics and statistical mechanics, but the link is established through model-specific choices—coarse-graining, ensemble definitions, and acceptance of certain limit theorems. These are framework decisions. We are not stepping outside; we are building a tighter frame.

Or take a DH example: a term's declining frequency across an eighteenth-century corpus might persist across tokenisation methods, statistical measures, and overlapping datasets. This robustness makes it valuable, but it remains indexical—true relative to the admissible set of methods and sources, not “the eighteenth-century mind as it really was.”

6.4 Two Readings of ISR

There are two coherent readings of ISR. Policy-ISR: treat invariants as our best available guide and pursue them aggressively; realism is a stance that motivates method. Result-ISR: treat invariants as sufficient evidence that we have learned the structure of reality “as it is.” My argument accepts Policy-ISR and builds it into IxSR. It rejects Result-ISR under current conditions because the missing truth test makes the realist conclusion underdetermined by the very LoAs that license it. Read as policy, ISR guides method and IxSR is its operationalisation. Read as result, ISR overreaches under present tests; IxSR blocks that slide.

6.5 Indexical Structural Realism

This leads to what I call *Indexical Structural Realism* (IxSR)⁵: progress about stable patterns, but always relative to admissible frameworks—those that are independent in their data, representation, or metric; linked by explicit, published translations; and candid about the norms they use to score success. This is valuable, but it is weaker than claiming to know the structure of reality “as it is in itself,” without mediation.

7. Additional Objections and Replies

The Structure-Dependence Problem, as formulated above, directly targets ISR’s attempt to combine LoA-bound epistemology with a robust realism about structure. It anticipates Floridi’s likely responses, such as appealing to the special ontological status of informational structures or pointing to invariants across frameworks. Yet the problem also invites broader philosophical pushback. Here I address five lines of objection likely to arise in peer review or interdisciplinary discussion.

7.1. Pragmatist Defence: “If It Works, It’s Knowledge”

A pragmatist might concede the LoA-bound nature of all access but deny that this undermines ISR’s realism. On this view, the distinction between “knowing via a structure” and “knowing reality as it is” collapses into practical success: if a claim generates reliable predictions, solves problems, and coordinates action effectively, then it counts as knowledge.

Reply:

Pragmatic success is not in dispute; LoA-bound methods can be extraordinarily effective. The point of the Structure-Dependence Problem is not that LoA-bound systems fail in practice, but that their success does not license the stronger metaphysical claim that we have captured reality’s structure *as it is in itself*. Practical success may be consistent with many incompatible underlying structures. Without a non-LoA-bound correctness test, we cannot justify the leap from “this works” to “this is reality’s structure.”

7.2. Deflationary Truth: “Truth Just Means What Our Best Framework Endorses”

A deflationist or conceptual relativist could argue that truth needs no correspondence beyond a framework. On this view, there is no “reality as it is” to compare to; truth is simply what our most coherent, well-supported LoA delivers. The Structure-Dependence Problem dissolves because the idea of a non-LoA-bound truth test is a category error.

Reply:

If truth is defined entirely within an LoA, then ISR’s claim to realism “as far as we can tell” is misleading. The position collapses into **internal realism** — Hilary Putnam’s view that truth is idealised rational acceptability within a conceptual scheme. This is a coherent philosophical position, but it is not what Floridi claims to be offering. ISR is meant to be a form of realism, not a deflationary or purely internalist account. If Floridi accepts deflationism, the Structure-Dependence Problem disappears — but so does the realist ambition of ISR.

7.3. Radical AI Epistemology: “AI Could Create New Forms of Knowing”

A more speculative objection looks to AI. Perhaps advanced AI systems could generate forms of knowledge that do not require human-style correctness norms or LoAs. On this view, the Structure-Dependence Problem is anthropocentric; it may apply to us, but not to AI epistemic agents.

Reply:

Any claim that an AI system “knows” presupposes some way of distinguishing correct from incorrect outputs. Whether that norm is designed by humans, learned through self-modification, or emergent from interaction with the environment, it still functions as a **framework**: it fixes observables, inference patterns, and standards of correctness. Even a radically alien intelligence would operate within some such structure, or else its outputs could not meaningfully be called “knowledge” rather than random behaviour. The Structure-Dependence Problem is not tied to human cognition; it applies to any system for which “knowing” implies the ability to separate correctness from error.

7.4. Interventionist Realism: “If You Can Manipulate It, It’s Real”

A related but distinct objection comes from the interventionist tradition in philosophy of science, associated with Ian Hacking’s maxim: “If you can spray them, they are real.” On this view, the capacity to successfully intervene in the world using a structural model provides a form of verification that is not merely representational. Building a particle accelerator based on quantum field theory or a CRISPR gene-editing tool grounded in molecular biology, and having them work as predicted, seems like the world “pushing back” in a way that confirms we have correctly captured its structure.

Reply:

Intervention unquestionably deepens our engagement with the world, but it does not escape framework dependence. The very criteria for what counts as a “successful intervention” are themselves defined within a framework: our LoA determines the operational goals, the acceptable error margins, the

measurement protocols, and the statistical thresholds for confirmation. The “push back” we register is mediated by instruments, interpretive models, and norms of success — all part of the same structured interface. Intervention can reinforce our confidence within a framework, but it cannot, by itself, provide the non-LoA-bound truth test that ISR’s realism requires.

7.5. Misinterpretation Charge (Pro-Floridi)

“You have built a straw man. ISR already bakes mediation into LoAs; invariants across well-constructed LoAs are exactly the non-arbitrary traction you seek. IxSR just restates ISR with weaker rhetoric.”

Reply:

This objection is directly addressed by the distinction between **Policy-ISR** and **Result-ISR** (defined in §VI), which is designed precisely to avoid this straw man. I argue that while Floridi’s work can be read as a defensible policy (which IxSR endorses and formalizes), it is often presented with the force of a metaphysical result. My critique targets this slide from policy to result, which the Structure-Dependence Problem makes untenable.

Therefore, the dispute is not about mediation or the value of invariants, both of which I grant. It is about the **status of the inference** from those invariants to realism. Invariants scored by LoA-given norms show robustness within an admissible family; to say they show correspondence beyond that family is to add a claim that cannot be certified by the same norms without circularity. IxSR is not a rebrand; it is a restriction with method. It (i) confines warrant to an index, (ii) imposes independence and translation requirements, and (iii) forces claim-scope statements to prevent the very slide from policy to metaphysics at issue. If ISR is read as policy, we converge; if read as result, we disagree for these principled reasons.

7.6 Why These Objections Matter

These objections illustrate the main strategies for trying to escape the Structure-Dependence Problem. Some attempt to redefine knowledge pragmatically or in terms of intervention — arguing that if a model works in practice, or allows us to manipulate the world successfully, then it must be latching onto reality. Others redefine truth in deflationary or coherentist terms, so that it no longer requires a correspondence beyond our frameworks. Still others imagine radically different epistemic agents, such as advanced AI systems, that might transcend our current cognitive limits. Addressing each of these routes reinforces the point that the problem is not parochial to a specific philosophy of information or to human cognition. It is a general challenge for any realism that combines structure-dependent access with the claim to capture reality “as it is in itself.”

These objections and replies clarify the evidential scope of invariants and the limits of the claims they can support. With that scope in view, it becomes essential to ask what Floridi’s optimism about the structure of reality being “learnable” actually amounts to, and how the Structure-Dependence Problem bears on different ways of understanding that term.

8. What ‘Learnable’ Can Mean

In assessing Floridi’s optimism about the world’s informational structure being “learnable,” it matters what we take that word to mean. There are at least three distinct readings in play, and the Structure-Dependence Problem bears differently on each.

On the strongest reading, “learnable” means the full capture of mind-independent reality—the idea that a method could yield the exact structure of the world as it is in itself. This is the realist ambition that gives ISR much of its appeal, but it is also the most vulnerable to the Structure-Dependence Problem. To vindicate it, one would need a correctness norm that does not originate in any framework and could be used to verify outputs without reintroducing mediation. As argued above, no such norm is available under present epistemic conditions. On this reading, the dilemma closes the door entirely: without a non-LoA truth test, the strongest form of learnability is blocked.

A more moderate, Peircean reading treats learnability as asymptotic convergence under ideal inquiry. Here, the claim is not that we possess the final structure now, but that our methods could in principle approach it over time as evidence accumulates and frameworks are refined. The Structure-Dependence Problem does not refute this possibility outright; it simply shows that, at present, ISR can only be justified as a research policy rather than a proven fact. Floridi’s optimism can survive on this reading, but its force changes: the invariants we observe become provisional signs to keep looking, not proof that we have already arrived.

The weakest reading treats learnability as convergence on structures that remain stable across a set of admissible frameworks. This is the reading most consistent with the methodological core of structural realism and the one fully compatible with IxSR. On this view, to say that the world’s structure is “learnable” is to say that we can build a resilient, interconnected web of frameworks in which certain patterns persist—patterns robust enough to guide practice, coordinate inquiry, and justify confidence within the bounds of those frameworks. The Structure-Dependence Problem leaves this reading untouched, but it also strips away any claim that such stability is evidence of correspondence with reality as it is in itself.

Seen this way, “learnable” is not a single thesis but a spectrum. The strong reading fails under present conditions; the Peircean reading reframes ISR as an aspirational policy; and the invariantist reading collapses ISR into IxSR, narrowing its epistemic ambition but preserving its methodological value. The task, then, is not to decide whether reality’s structure is learnable in the abstract, but to be explicit about which sense of the term we mean—and what kind of warrant each sense can support. Clarifying these readings of “learnable” sets the stage for a measured middle ground: a stance that acknowledges our present epistemic limits while leaving open the possibility of future breakthroughs in access, modelling, or verification.

This is where misinterpretation typically arises. Proponents hear my view as denying learnability altogether. I deny only the strong reading—that invariants certify mind-independent structure now. I accept the Peircean and invariantist readings as policy and method: pursue invariants, increase

independence, publish translations, and index every claim. IxSR codifies that discipline and names the limit it leaves in place.

9. Bridging Present Limits and Future Possibility — A Measured Middle Ground

This measured middle ground, emerging from the three readings of “learnable” in §VIII, recognises that while our current epistemic condition blocks the strongest form, it does not commit us to permanent agnosticism. We can value Floridi’s optimism as a research policy—pursuing greater access and more robust invariants—while resisting the temptation to treat that optimism as an achieved fact.

Two questions must be kept distinct: *Is it possible in principle to know mind-independent reality?* and *Are we justified in claiming that we can?* My critique addresses the latter. I do not assert that our present limits are fixed for all time, nor that future breakthroughs in cognition or modelling are impossible.

Floridi’s optimism belongs to a long tradition of intellectual expansion. As a research policy, it can motivate innovation, the creation of new methods, and the pursuit of ambitious theories. There is value in treating greater access as possible until proven otherwise.

9.1 Thought experiment — the non-LoA truth test

Imagine an oracle that prints “REAL.” Without a framework to test, calibrate, or even interpret the output, we cannot rank success over failure. Add calibration — and you have added a framework. Either way, the miracle dissolves.

But present conditions do not justify certainty. Without a truth-test independent of LoA mediation, claims to have captured the structure of mind-independent reality risk conceptual overreach. The Structure-Dependence Problem means that even a radical epistemic system would have to show how it could generate knowledge without reinstating a framework of access.

None of this blocks Floridi’s optimism as a research policy. It blocks only the slide from policy to result under present tests. If a LoA-external correctness norm becomes available, Result-ISR would gain its missing warrant. Until then, responsible realism is **indexed**: say exactly which LoAs and translations support a claim, and resist reading cross-LoA stability as correspondence without further argument.

10. Conclusion — Hope, Caution, and the Shape of Knowledge

This paper has argued that while Floridi’s informational structural realism is an ambitious and appealing account of how we might know reality’s structure, it faces the Structure-Dependence Problem: any method that counts as knowledge must operate within a framework, leaving no way to verify claims about “reality as it is in itself” without reintroducing framework-dependence. As the truth fork in §III made clear: if you deflate truth, you lose the realism Floridi advertises; if you keep correspondence, you need a non-LoA truth test we do not have.

The Structure-Dependence Problem adds a further challenge: knowledge, as we understand it, may be inextricable from such structures. If a future method lacks them, it is not knowledge; if it retains them, it is still structure-bound. Cross-LoA invariants offer robustness, but only indexical progress — still tied to admissible frameworks.

For the digital humanities, this means that AI and computational tools cannot escape the epistemic boundaries of their own design frameworks. Cross-tool invariants and robust outputs are valuable, but they are still indexical — progress relative to admissible frameworks, not unmediated knowledge of cultural or historical reality. Recognising the Structure-Dependence Problem in digital contexts is not a call to abandon these tools, but to integrate epistemic humility into their use.

10.1 Operationalising IxSR in Digital Humanities and AI

In short, under present conditions ISR is best read as a research policy rather than a metaphysical result: invariants license use, not correspondence. If a LoA-external truth test emerges, that verdict can change; until then, responsible realism is indexed.

As a research policy, IxSR licenses a constructive and rigorous program within the acknowledged limits of framework-bound knowledge. In practice, this means:

10.2. Design for multi-LoA coverage

Designing for multi-LoA coverage means building systems and methods that can operate under multiple representational schemes, parameter settings, or ontologies. In computational literary studies, for example, a project might run both a probabilistic topic model and a neural embedding model on the same corpus, comparing the stability of themes across these very different LoAs. In cultural heritage, a database could be structured to support both a simple object-based schema and a more complex event-based ontology. By ensuring that research questions are tested across more than one LoA, scholars increase the chances of identifying invariants that are robust within a diverse set of admissible frameworks.

10.3. Prioritize interoperability

Prioritizing interoperability means actively designing systems that can “talk” to each other. In the cultural heritage sector, this could involve funding projects that create robust mappings between different metadata standards, such as Dublin Core and CIDOC-CRM. In AI research, it might mean developing methods to align the vector spaces of different language models so that outputs can be meaningfully compared. The goal is not to impose a single “master framework,” but to build the connective tissue that allows invariants to be tested across a diverse ecosystem of LoAs, thereby strengthening the entire web of knowledge.

10.4. Systematically search for invariants

Searching for invariants should be a deliberate and documented research practice, not an incidental byproduct. In a DH context, this might involve running the same query or analytical workflow across multiple datasets, tools, and ontologies, then cataloguing which patterns persist and which disappear. In AI, a team might compare the outputs of different models trained on different corpora, seeking patterns that are stable despite differences in architecture or training data. This approach reframes “results” not as definitive truths, but as points of convergence within a network of diverse frameworks.

10.5. Audit framework-dependence

Auditing framework-dependence demands a new standard of methodological transparency. A published paper using topic modelling, for instance, would not simply present the final topics as a discovery. It would include an appendix or link to a repository detailing the specific LDA library used, the range of topic numbers tested, the exact stop-word list applied, and the justification for the final parameter choices. In HTR work, an audit would detail the ground truth transcription guidelines, the segmentation algorithms, and the criteria for accuracy. This makes the LoA explicit and allows the research to be understood not as a statement about “the text itself,” but as a finding within a specific, reproducible analytical framework.

10.6. Interpret invariants with humility

Finally, interpreting invariants with humility means resisting the temptation to frame robustness across LoAs as access to reality “as it is in itself.” Instead, invariants should be presented as warrant for use and trust *within practice* — a sign of methodological strength, not metaphysical finality. For example, if a thematic pattern in a corpus persists across both topic modelling and word embeddings, that persistence justifies treating it as a reliable feature for further study, teaching, or exhibition. But its ontological status beyond the frameworks that produced it remains an open question. This humility safeguards against overreach while keeping the focus on building an increasingly resilient and interconnected body of knowledge.

10.7. Say what backs each claim (“the index”)

Under every substantive result, add a short **Index** line that tells readers three things in plain English:

- (1) **Which frameworks** you used (the LoAs),
- (2) **How you matched** results across them (the bridge), and
- (3) **What counted as success** (the rule or threshold you used).

Template you can drop in:

Index: Frameworks — [Framework A; Framework B]. Bridge — we treated two outputs as “the same” when [simple rule]. Success rule — a result counted as strong when [plain rule or threshold].

Example (literary NLP):

Index: Frameworks — topic model; embedding model. Bridge — we called them “the same theme” when most top words overlapped and two reviewers agreed. Success rule — topics had to meet our

coherence threshold and clusters had to meet our similarity threshold.

Example (handwritten text recognition):

Index: Frameworks — model trained on Guideline A; model trained on Guideline B. Bridge — outputs were matched after standardising punctuation. Success rule — a line counted as accurate when its character-error rate was five percent or lower.

Example (museum data):

Index: Frameworks — source collection schema; mapped CIDOC-CRM view. Bridge — creators matched by authority IDs; dates normalised. Success rule — queries passed when validation checks succeeded and curators spot-checked them.

10.8. Say how independent your checks were, and what was lost in translation

Add one sentence that covers two points:

- (1) **Independence** — did your cross-checks differ in **data**, **representation**, or **metrics**? (Different corpus? Different modelling family? Different scoring rule?)
- (2) **Translation debt** — what got simplified, blurred, or could be wrong when you mapped one framework to another?

Template you can drop in:

Independence: [Brief note on whether data/representation/metrics were the same or different.]

Translation debt: [Brief note on any ambiguous mappings or known losses.]

Example (literary NLP):

Independence: Same texts; different modelling families and different scoring rules. **Translation debt:** The “factory/works” synonym mapping is ambiguous; rare terms can drop out when we align vocabularies.

Example (handwritten text recognition):

Independence: Different training sets; same model architecture and same accuracy measure.

Translation debt: Differences in how ligatures are transcribed may inflate error when results are compared.

Example (museum data):

Independence: Same objects; different representations (object list vs event-based ontology) and different validation methods. **Translation debt:** Mapping movements to events can lose some early-period detail.

In short, ISR’s mediation is not in dispute; the inference is. Cross-LoA invariants justify use and coordination **inside** an admissible family; they do not, by themselves, license correspondence claims **outside** it. **IxSR** turns that limit into method—independence, translations, robustness sweeps, and indexed claims—so we keep the traction invariants provide without overstating what they prove.

We can hope to know more. But hope, however motivating, is not yet knowledge.

References

- Drucker, J. (2011). Humanities approaches to graphical display. *Digital Humanities Quarterly*, 5(1).
- Floridi, L. (2011). *The philosophy of information*. Oxford University Press.
- Floridi, L. (2019). *The logic of information: A theory of philosophy as conceptual design*. Oxford University Press.
- Humphreys, P. (2004). *Extending ourselves: Computational science, empiricism, and scientific method*. Oxford University Press.
- Levins, R. (1966). The strategy of model building in population biology. *American Scientist*, 54(4), 421–431.
- Newman, M. H. A. (1928). Mr. Russell's causal theory of perception. *Mind*, 37(148), 368–374.
- Putnam, H. (1981). *Reason, truth and history*. Cambridge University Press.
- Rockwell, G., & Sinclair, S. (2016). *Hermeneutica: Computer-assisted interpretation in the humanities*. MIT Press.
- Steel, D. (2008). *Across the boundaries: Extrapolation in biology and social science*. Oxford University Press.
- Underwood, T. (2019). *Distant horizons: Digital evidence and literary change*. University of Chicago Press.
- van Fraassen, B. C. (1980). *The scientific image*. Oxford University Press.
- Weisberg, M. (2006). Robustness analysis. *Philosophy of Science*, 73(5), 730–742.
<https://doi.org/10.1086/507532>
- Weisberg, M. (2013). *Simulation and similarity: Using models to understand the world*. Oxford University Press.
- Wimsatt, W. C. (1981). Robustness, reliability, and overdetermination. *PSA: Proceedings of the Biennial Meeting of the Philosophy of Science Association*, 1981(2), 124–163.
- Worrall, J. (1989). Structural realism: The best of both worlds? *Dialectica*, 43(1–2), 99–124.
<https://doi.org/10.1111/j.1746-8361.1989.tb00933.x>

Notes

1. By ‘role-preserving’ I mean preserving measurement or causal roles across translations, not merely point-wise relabelling that could be satisfied by arbitrary set-theoretic constructions. ↩
2. I use “invariant” non-trivially: preserved under translations that keep measurement scales, causal/explanatory roles, or intervention targets—not under arbitrary re-labelling (cf. Newman 1928). ↩
3. For clarity: I use “index” to denote the tuple $\langle \text{LoAs}, \text{translations}, \text{norms} \rangle$ that underwrites a given claim. ↩
4. This parallels robustness analysis in the philosophy of science, where agreement across independent models or methods (Levins 1966; Weisberg 2006) increases confidence without assuming any one is perfectly accurate. ↩
5. Related in spirit to positions that treat scientific models as perspectival or partial (e.g., Giere 2006; van Fraassen 2008), but tied here to a formal admissibility condition on frameworks. ↩

Recommended citation

Whittingham, Nathan G. 2026. *The Axiom of Hope: Indexical Realism and the Structure-Dependence Problem*. Preprint, v1.0, 27 January. <https://nathanwhittingham.com/papers/the-axiom-of-hope/>