

Interpretability Engine Optimization

Methodology · Version 1.0

Juan Carlos López Castaño

ORCID: 0009-0008-7471-0650

NSOLVIA Research

Version 1.0 · 5 September 2026

DOI: 10.5281/zenodo.22416784

License: CC BY 4.0

Descriptor: The UNDERSTAND Layer

Suggested citation: López Castaño, J. C. (2026). *Interpretability Engine Optimization: Methodology*. NSOLVIA Research. Version 1.0. DOI: 10.5281/zenodo.22416784

Companion to: *Interpretability Engine Optimization (IEO): The Missing Optimization Layer in AI Commerce*. DOI: [10.5281/zenodo.22104280](https://doi.org/10.5281/zenodo.22104280)

Terminology: nsolvias.com/ieo/glossary

Contents

	Preface	3
1	The objective	3
2	Definitions	4
3	The five levels	5
4	Principles	6
5	The four failure modes	8
6	The cycle	9
7	The decision context as the unit of work	10
8	Measurement	11
9	Evaluation and falsifiability	12
10	Domain profiles	13
11	How to build a domain profile	15
12	Who practices this	16
13	The research agenda	18
14	Relation to adjacent practices	18
15	How this methodology evolves	19
16	Limits	20
17	Terminology	21
18	On openness	21
	Version history	21

Preface

The founding paper proposes the discipline and reports its first evidence. This document states **how the discipline is practiced** — its objective, its principles, its unit of work, its failure modes, its cycle, its measurement, and the way it is instantiated in a domain.

It is written so that a practitioner with no relationship to NSOLVIA can read it and practice. That is the test of whether Interpretability Engine Optimization is a discipline rather than a product, and this document exists to pass it.

What this document is not. It does not describe how NSOLVIA's own systems implement the discipline. It specifies the objective and the evaluation discipline; it does not prescribe an implementation. That separation is deliberate and it is stated in the first principle below.

Scope. This methodology is **domain-general**. The objective it defines, the failure modes it names, the cycle it describes and the tests it specifies apply wherever a machine must reason or act on business information — commerce, credit, regulatory compliance, clinical information, legal documents, public infrastructure, and domains not yet named.

Status of the evidence. The discipline has been **measured in one domain — commerce** — with one instrument and one versioned corpus. That distinction is maintained throughout and it is not incidental: **where this document describes the discipline, it applies to any domain; where it reports evidence, the evidence is from commerce.** Every other domain named here is named as a candidate, not as an established profile, and the document says so wherever it matters.

1. The objective

Interpretability Engine Optimization (IEO) is the practice of making source information machine-interpretable for a declared decision context before downstream reasoning, recommendation, comparison, decision, or action occurs.

That is the general statement. In commerce — the first domain in which the discipline was instantiated and measured — it takes the form stated in the founding paper: *the practice of making a business's offers, entities, attributes, policies, and source facts machine-interpretable before downstream retrieval, recommendation, comparison, or action occurs.* The commerce formulation is a domain profile of the general one, not a competing definition.

Stated operationally: there is information in a domain; its meaning is resolved with the level of control the domain requires; a representation is produced that a machine can interpret correctly for its task; that representation feeds reasoning, decision or action.

The objective is interpretation, not visibility. IEO does not try to make a machine prefer an offer. It tries to ensure that the machine can determine correctly whether the offer satisfies what was asked — when it does and when it does not. It corrects both errors: the eligible offer that went unrecognized, and the ineligible one recommended on ambiguous data.

If an offer is genuinely unsuited to a task, IEO does not make it suited. It makes that fact resolvable.

The first principle

IEO specifies the objective and the evaluation discipline. It does not prescribe the implementation stack.

Normalization, taxonomy resolution, enrichment, grounding, structured data, ontologies, knowledge graphs, retrieval-augmented generation, semantic layers — all of these predate IEO and all of them may serve it. None of them defines it. An organization may practice IEO with none of them, with some of them, or with technology that does not yet exist. What makes the practice IEO is the objective it optimizes and the discipline by which it verifies the result.

2. Definitions

Interpretability. The degree to which a machine can correctly resolve, from a representation, what it needs in order to reason or act — without inference the source does not support.

Interpretability is not readability. A machine reads a representation when its characters parse. It interprets it when it can resolve identity, taxonomy, attributes, constraints and context well enough to act correctly. A representation can be perfectly readable and poorly interpretable.

Source representation. The information as it exists before IEO acts on it: pages, records, documents, databases, feeds, in whatever form the domain produces them.

Interpretable representation. The output of the practice: a representation whose meaning is verified, semantically resolved, and adequate to the decision context it serves. It may take any form — a structured record, a graph, an API response, a document with declared constraints — and the form is not part of the definition.

Decision context. The task for which a machine must interpret the representation. The same source can require different resolutions for different tasks. A credit applicant's data must be interpreted differently for a credit card, a car loan and a mortgage; a compliance report must be interpreted differently for poultry housing and for cattle; a product must be interpreted differently for a purchase recommendation and for a returns eligibility check. The decision context determines what "correctly interpreted" means.

Interpretability Engine. Any machine system that consumes business information and must form a usable interpretation of it — search engines, answer engines, generative systems, recommendation systems, commerce platforms, autonomous agents, decision engines.

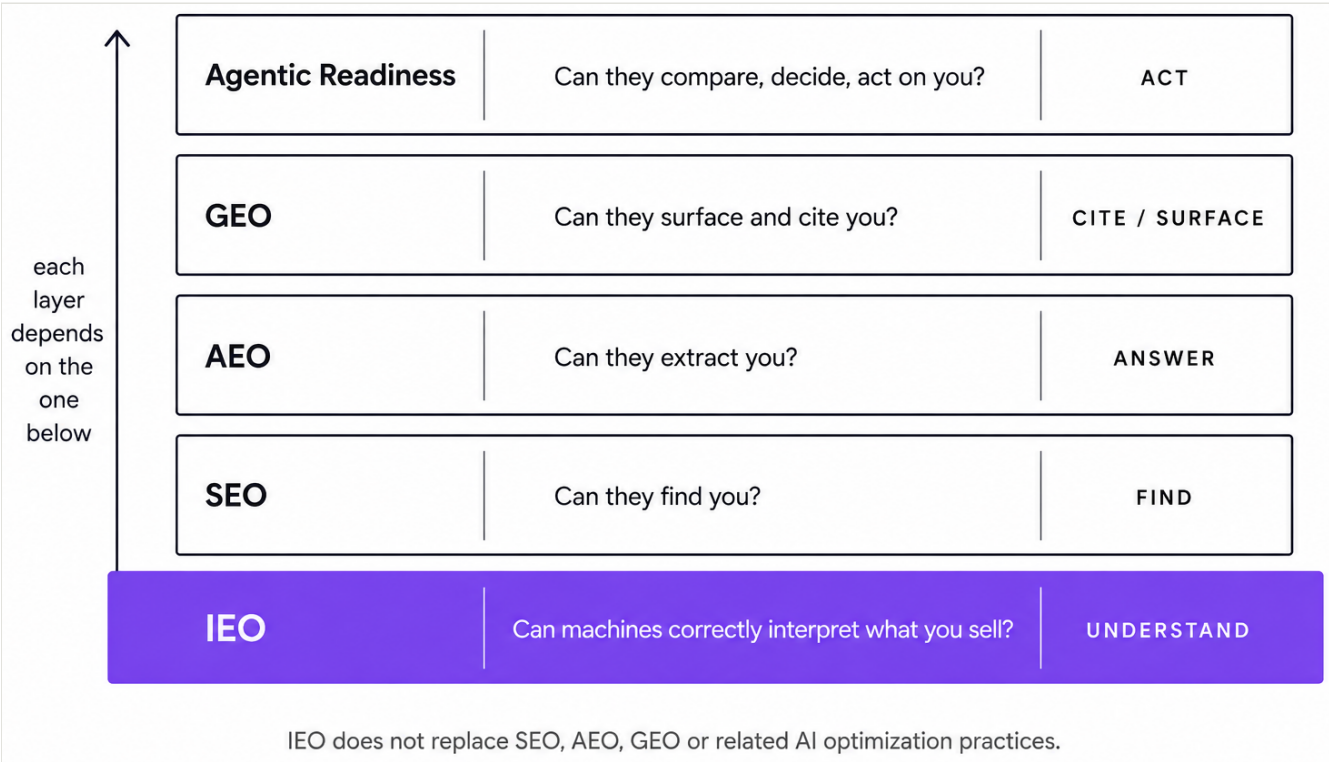


Figure 1. The IEO ladder: interpretability as the layer preceding discovery, extraction, citation and action.

3. The five levels

Confusion about IEO almost always comes from mixing levels. Kept separate, they are clean.

Level	What it is	Example
1. IEO	The general principle and this methodology. Domain-agnostic.	This document
2. Domain Profile	What interpretability means in a specific domain: which objects, which resolutions, which evidence, which controls.	Commerce · credit assessment · regulatory compliance
3. Implementation	How an organization builds a system that practices the discipline. Proprietary or open, as the organization chooses.	Any organization's engine, in any domain
4. Instrument	How interpretability is measured in that domain. Any organization may build its own; multiple instruments may coexist and are expected to disagree in examinable ways.	A catalog readiness audit · an interpretability check for a given decision context
5. Application	The concrete problem a given deployment solves.	Preparing a catalog for AI shopping surfaces · resolving what a risk indicator means for a specific credit product

Only Level 1 is general. Levels 2 through 5 are domain-specific by construction — the profile, the implementation, the instrument and the application all change with the domain. The examples above are drawn from more than one domain deliberately: the methodology is not a commerce methodology with other uses appended. It is a general methodology whose evidence, so far, is from commerce. Section 10 lists candidate domains; Section 11 states how a profile is built.

A domain expert provides the domain truth. An IEO practitioner defines how that knowledge becomes explicit, resolvable, verifiable and testable by machines. An engineer implements. An agent consumes. No single party needs to occupy all five levels, and the discipline does not require any one organization to.

The methodology below lives at Level 1. Section 10 describes Level 2 and how to build one. Levels 3 through 5 belong to whoever practices.

4. Principles

4.1 — Stack-agnostic. Stated above as the first principle. IEO specifies the objective and the evaluation discipline; it does not prescribe the implementation.

4.2 — Verified, never invented. Every element of an interpretable representation traces to evidence the source supports. **IEO does not prohibit inference. It prohibits unsupported inference.** Resolving "100% cotton" to `material = cotton`, or a product to its taxonomy category, is a semantic transformation — and it is legitimate because it is entailed by authoritative evidence. What is not legitimate is a value the evidence does not support. Where evidence is insufficient, the element remains explicitly unresolved — an *honest empty* — rather than guessed.

Every assertion therefore carries a declared state: **explicit** (stated in the source), **derived** (entailed from the source by a stated rule or reference), **inferred with confidence** (probabilistic, with its confidence and basis declared), **unresolved** (evidence insufficient), or **conflicting** (sources disagree, not yet resolved). Probabilistic inference may be represented when the interpretability contract permits it, but its inferential status, basis and confidence must remain explicit; **it is never promoted to verified fact.** A representation that does not declare these states can no longer be evaluated, because resolved values and guesses become indistinguishable.

4.3 — The mirror principle. Human-facing and machine-facing layers state the same facts because they derive from the same source. The equivalence is **factual, not literal**: a machine-facing representation may carry taxonomy identifiers, canonical references, normalized units, confidence values and provenance pointers that no human reads on a page — but **every factual assertion in it must be traceable to human-verifiable evidence or to an authoritative reference.** This makes the practice auditable by construction and rules out presenting different truths to different readers.

4.4 — Reproducibility as the domain requires. *Interpretation should be made as reproducible, bounded and auditable as the risk profile of the domain requires.* Components may be probabilistic; what matters is that rules, provenance, confidence, thresholds, guardrails, honest empties and human review are applied in proportion to the cost of error. In commerce a misinterpretation costs a return. In regulated domains it may cost a certification, a diagnosis, or a legal outcome. The same discipline; a rising floor of control.

4.5 — IEO prepares interpretation; it does not make the decision. IEO resolves what a representation means for a task — that a risk flag means *high mortgage default risk under a stated definition, with provenance and date*. The system that decides what to do with that meaning is a different system, governed by its own discipline. IEO does not replace risk models, clinical decision support, underwriting engines, or regulatory decision systems. It feeds them a representation they can interpret correctly.

4.6 — Same objective, domain-specific mechanisms. *IEO applies the same optimization objective across domains; each domain defines the representations, evidence, constraints and controls required for correct interpretation.* Commerce may resolve against taxonomies and attributes. Banking may require ontologies and policy rules. Regulatory compliance may require regulations, evidence chains, temporal validity and jurisdiction. The principle is common; the mechanisms are not.

4.7 — Measurement before claims. Nothing is asserted about the effect of an intervention that has not been measured, and every figure is labeled by what it is: a measurement, a projection, or an observed post-deployment result. Section 8 defines the terms.

5. The four failure modes

Semantic failures fall into four modes. A separate representational condition may also prevent interpretation. Diagnosing which applies determines the remedy.

Four semantic failure modes.

5.1 — Ambiguity. A value exists but does not resolve to a single meaning. A category set to an internal label. A product name that states a brand and a code but not what the item is. A field `risk = high` with no indication of which risk. *Remedy: resolution against a controlled reference.*

5.2 — Conflict. The same fact is stated differently in different places, or two elements contradict each other. A return window of thirty days on one page and fifteen on another. A material declared as cotton in one field and polyester in the description. *Remedy: resolution with the source of truth, then a single statement. Never a silent choice.*

5.3 — Absence. The information does not exist in the source. No identifier anywhere; no stated constraint; no provenance. *Remedy: an honest empty, with guidance on what it would take to resolve. Never inference.*

5.4 — Unsupported inference. A value exists but was not derived from the source — it was guessed, generated, or carried over from something similar. This is the failure mode that looks like success: the field is populated, the validator passes, and the value is unverifiable. *Remedy: retraction to honest empty, then resolution from evidence.*

The fourth mode is the most expensive, because it is invisible until it fails downstream.

A separate condition, distinct from the four failure modes: representation deficiency. The evidence exists — in prose, in a paragraph, legible to any human — and is not expressed in a form a machine can resolve for the decision context. This is not a semantic failure of the fact itself; it is a deficiency in how existing evidence is represented. It is diagnosed separately because its remedy is different — structuring, not resolution — and it is the most common condition in most source representations.

6. The cycle

The practice is a loop, not a project.

Measure → Find → Remediate → Validate → Re-measure

Measure. Establish the interpretability baseline of the source representation, on a versioned instrument, for the decision context in question.

Find. Diagnose each failure — ambiguity, conflict, absence, unsupported inference, unstructured presence — element by element.

Remediate. Resolve what can be resolved from evidence. Leave honest empties where it cannot. Structure what exists as prose. Never invent.

Validate. Confirm that every remediated element traces to the source or to a confirmed fact. Confirm that the human-facing and machine-facing layers agree. Confirm that nothing unsupported was introduced.

Re-measure. Run the same instrument on the remediated representation of the same underlying object, under the same decision context and evaluation conditions. The difference between the two measurements is the observed lift. If the instrument or its scale changed between runs, the measurements are not comparable and the earlier one is discarded as a baseline rather than rescaled.

The cycle repeats, because sources change and the machines reading them change. A representation interpreted correctly today is not guaranteed to be interpreted correctly after the source is edited or the consuming system is updated. Maintaining interpretability is part of the practice, not an afterthought to it.

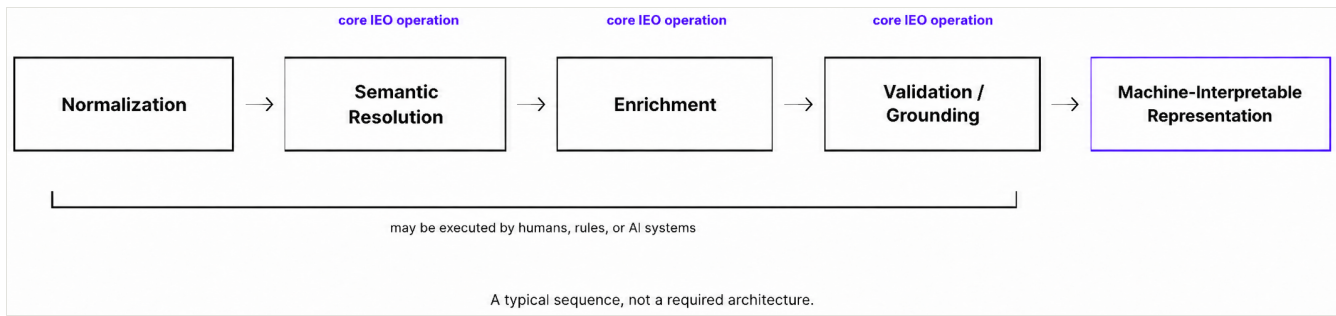


Figure 2. *The IEO pipeline. A typical sequence, not a required architecture.*

7. The decision context as the unit of work

The unit of IEO is not the record, the page, or the field. It is **the interpretation of a representation for a decision context**.

This has a practical consequence: interpretability is not a single property of a source. A source can be well interpretable for one task and poorly interpretable for another. The practitioner's first question is therefore not *"is this representation interpretable?"* but *"interpretable for what?"*

Defining the decision context means stating:

- **What the machine must resolve** to perform the task — which entities, attributes, constraints, relationships, rules.
- **What counts as correct** for each of those — the reference against which a resolution is judged.
- **What the cost of error is** — which sets the level of control the domain requires under principle 4.4.

Two decision contexts that share a source may require different profiles of resolution. That is not a complication of the discipline; it is its content.

7.1 The unit of evaluation

A decision context is evaluated through its **resolution requirements** — the specific elements a machine must correctly resolve to perform the task. This is the unit that makes two practitioners' instruments comparable.

Decision context → Resolution requirements → Gold truth → Machine resolution → Evaluation

Each resolution requirement declares six things: the **expected value or accepted set** · the **reference** against which it is judged · the **provenance** required · the **resolution state** (explicit, derived, inferred, unresolved, conflicting) · its **criticality** · and its **cost of error**.

Example, commerce. Decision context: *recommend a running shoe for a vegan customer under \$150*. Resolution requirements: product identity · category · price · material composition · vegan constraint · availability. Getting the color wrong is a low-cost miss. Getting the material wrong violates the constraint the customer stated. The requirements are not equal, and the contract says so.

7.2 The interpretability contract

For each decision context, a domain profile states an **interpretability contract**: what must be resolved · against what truth · with what evidence · what uncertainty is tolerated · what must remain unresolved · and **which errors are inadmissible**.

The contract is what turns the discipline from a philosophy into an auditable practice. An instrument measures against a contract. A remediation is validated against a contract. A third party can check a contract without trusting anyone.

8. Measurement

Three terms, and they carry different evidential weight.

Interpretability Baseline. The observed current state, measured from the source representation on a versioned instrument. A baseline is a measurement: run twice under identical conditions on an unchanged source, it returns the same value **within a tolerance the instrument declares**. A deterministic instrument declares zero tolerance; an instrument with probabilistic components declares its variance. Either is acceptable. An instrument that does not declare which it is cannot produce a baseline.

Projected Interpretability Lift. The estimated improvement represented by a remediated version, produced before deployment. A projection is an estimate: the remediated representation is generated, and generated representations can vary between runs even when the measurement of them does not.

Observed Interpretability Lift. The difference measured after remediation has been deployed and the representation re-evaluated on the same instrument. An observed result is the strongest evidence the discipline produces. It requires deployment, and it is typically smaller than the projection, because the projection estimates the ceiling if every remediation were applied and any deployment is a subset of it.

The comparability rule. Two measurements are comparable only if they share the same instrument version and the same dimensional maxima. A difference computed across a scale change is an artifact, not a lift.

On instruments. IEO does not prescribe a single score or a single instrument. Any tool that measures whether a machine can correctly resolve a representation's meaning for a task is measuring interpretability. A commercial readiness instrument may include interpretability measurements alongside downstream signals and should not be read as a universal IEO metric. The discipline invites independent instruments, and it expects them to disagree in ways that can be examined.

Known bias to declare. An instrument that registers whether a field is *populated* is not measuring whether its content is *interpretable*. Two representations with equally complete but unequally resolvable content can receive similar scores. This biases baselines upward, meaning reported failure is conservative. Know the bias direction of any instrument used.

9. Evaluation and falsifiability

Two tests. The first tells a practitioner whether an intervention belongs to the discipline. The second tells anyone whether the discipline's premise holds.

9.1 The Boundary Test

Did the intervention improve the machine's understanding of the source — or did it only improve how the source is formatted, distributed, discovered, or surfaced?

Practical form: *if the same information were delivered through a different channel tomorrow, would the improvement still remain?* Interpretability survives the channel. Positioning does not.

A representation may be valid structured data, perfectly normalized and widely distributed, and still be semantically ambiguous. IEO evaluates the ambiguity that remains.

9.2 The Counterfactual Test

The central claim — that a more interpretable representation produces better machine understanding of the same offer — is falsifiable.

Take one fully processed representation. Produce a controlled degraded version: category generalized, attributes removed, constraints deleted, everything else held constant. Give the same machine system the same questions and access to both. Measure attribute extraction accuracy, category resolution, constraint satisfaction, differentiation between near-identical items, hallucination rate and task completion — against a golden interpretation record defined in advance.

Protocol: same model and version, same configuration, same questions, same retrieval settings; representations evaluated separately and blinded; randomized order across multiple runs; gold labels fixed before results; confidence intervals reported.

If processed representations do not outperform degraded ones under these conditions, the premise fails. The discipline invites replication, and it means it.

A declared limit of the basic design. Degrading category, attributes and constraints together shows that removing useful information harms the model; it does not by itself separate the effect of semantic resolution from the effect of information availability or of format. An **ablation extension** — processed vs. source vs. degraded-category vs. degraded-attributes vs. degraded-constraints vs. a formatting-only control — would separate those effects and is listed in the research agenda.

10. Domain profiles

A **domain profile** states what interpretability means in a specific domain: which objects must be resolved, against which references, with what evidence, under what controls, for which decision contexts.

In operational terms, a domain profile is that domain's **standard operating procedure (SOP) for interpretation**: what must be resolved, against what, on what evidence, under what controls, and what is left explicitly unresolved. Practitioners in health, finance, manufacturing and public administration will recognize the form; what changes is that the procedure is written so that machines, not only people, can follow its output.

10.1 Commerce IEO — the first profile, and the only one with evidence

Commerce is the domain in which this discipline was measured, remediated and tested. Its profile is stated in the founding paper and summarized here.

Objects of interpretation: three.

- **Product** — what is being sold: identity, taxonomy, attributes, materials, use cases, functional intent, purchase signals, restrictions, and the differences between near-identical items.

- **Brand** — who stands behind it: identity, provenance, manufacturing facts, certifications, claims and the evidence supporting them.
- **Transaction** — the rules under which it can be bought: price, availability, shipping, returns, eligibility, exclusions.

References: official platform taxonomies; standardized attribute vocabularies; merchant-controlled or merchant-confirmed evidence, together with authoritative references where applicable.

Evidence standard: every value traces to what the merchant published or confirmed, or to an authoritative reference.

Typical decision contexts: recommendation, comparison, constraint-satisfying selection, purchase, returns eligibility.

Maturity: measured at the product level with a versioned instrument and a reproducible baseline; brand and transactional interpretation follow the same logic and are earlier in measurement maturity.

Artifacts: Semantic Product Record, Semantic Brand Record, Semantic Transaction Record.

10.2 Other domains — named as possibilities, not as profiles

The objective generalizes. The evidence does not yet.

The problem IEO names — information a machine can read but not correctly interpret for its task — exists wherever a machine will reason or act on business information. The following are stated as **candidate domains**, not as profiles. Each is written as an illustration of the same structure, and none is claimed as established work.

Candidate domain	What would be interpreted	Against what reference	Why the control level rises
Credit and financial products	Applicant attributes, risk indicators, product terms	Policy definitions, regulatory frameworks	The same indicator means different things for a card, a car loan and a mortgage; a misread flag is an incorrect decision about a person
Regulatory compliance	Inspection reports, standards, certifications	Regulations, jurisdiction, validity period	A facility standard for poultry is not the standard for cattle; a misinterpretation grants or denies a certification
Technical documentation	Procedures, dependencies, versions, constraints	Product versions, standards	A procedure applied to the wrong version fails silently

Three are listed; others — clinical information, legal documents, public infrastructure — follow the same structure and are named in the research agenda. Two observations follow from the table, and both are principles from Section 4 applied.

First, **the same source serves different decision contexts and requires different resolutions for each** — the credit example is the clearest case. Second, **IEO prepares interpretation; it does not make the decision**. In none of these domains does IEO replace the underwriting engine, the regulatory decision, the clinical judgment, the legal analysis or the engineering model. It resolves what the information means, with provenance, so that the system or person deciding is not deciding on ambiguity.

No profile for any of these domains is established by this document. Establishing one requires the work described in Section 11, done by people who hold the domain truth.

11. How to build a domain profile

A domain profile is built by a domain expert and an IEO practitioner together. Neither can do it alone.

Step 1 — Name the decision contexts. Which tasks will machines perform on this domain's information? For each, what must be resolved and what counts as correct?

Step 2 — Name the objects. What are the things a machine must interpret? Commerce has three; a domain may have more or fewer. Each object must be definable without reference to any implementation.

Step 3 — Name the references. Against what is a resolution judged? Taxonomies, controlled vocabularies, regulatory definitions, standards, jurisdiction-specific rules. If no reference exists, an element cannot be resolved arbitrarily: it remains an honest empty, or the domain profile formally establishes a local, verifiable reference — a controlled vocabulary or canonical list — and states that it did so.

Step 4 — Set the evidence standard. What must a value trace to before it is asserted? In commerce: the merchant's published pages. In regulated domains: typically a documented source with provenance, date and validity period.

Step 5 — Set the control level. Apply principle 4.4. What does a misinterpretation cost in this domain, and what reproducibility, bounding, auditability and human review does that cost require?

Step 6 — Build or adopt an instrument. How will interpretability be measured for this domain's decision contexts? The instrument must be versioned and its baseline reproducible before any lift is reported.

Step 7 — Establish the baseline, then run the cycle. Measure first. Everything after is the practice.

Step 8 — Declare the maturity. State plainly which objects and contexts have been measured and which have not. A profile with declared limits is a profile others can build on.

On the order. These eight steps describe what a domain profile must contain, not the sequence in which it must be discovered. In practice most profiles are built the other way around: someone solves a concrete interpretation problem, builds something that works, and only afterwards formalizes what was being measured and why. The commerce profile was developed that way — the instrument existed before the framework that describes it. **What matters is that all eight pieces exist and are explicit, not which one came first.**

A domain profile is open by construction: anyone may build one, and building one does not require permission from NSOLVIA or the use of any NSOLVIA tool.

12. Who practices this

A method is executed. A discipline is **practiced** — by people who hold a competence and a role.

12.1 What an IEO practitioner does

An IEO practitioner works between the people who hold a domain's truth and the systems that must act on it. The work is not writing content and it is not building infrastructure. It is making meaning explicit, verifiable and measurable.

A practitioner must be able to:

- **Define a decision context** — state what a machine must resolve for a given task, and what counts as correct for each element.
- **Establish a baseline** — measure the current interpretability of a source representation on a versioned instrument, and know what the instrument does and does not see.
- **Diagnose failure** — distinguish ambiguity from conflict, absence, unsupported inference and unstructured presence, because each has a different remedy.
- **Apply the boundary test** — tell whether a proposed intervention improves understanding or only formatting, distribution or presentation.
- **Decide what stays unresolved** — recognize when evidence is insufficient and leave an honest empty, with guidance, rather than fill it.
- **Set the control level** — judge what reproducibility, provenance, thresholds, guardrails and human review the domain's cost of error requires.
- **Report honestly** — separate measured, projected and observed, and never present one as another.
- **Re-measure** — run the cycle, and know when two measurements are not comparable.

Domain expertise is not on that list, and the omission is deliberate: **the practitioner does not supply the domain truth. The domain expert does.** The practitioner makes it explicit, resolvable and testable. Neither role can do the work alone, and confusing them is the most common way a profile goes wrong.

12.2 Where the work sits

The practice appears wherever an organization already owns the gap between what it knows and what its systems can use: data and catalog teams, information architecture, compliance and quality functions, domain informatics, and the agencies serving them. Many practitioners are already doing this work without a name for it.

12.3 On certification

There is no certification for this discipline, and this document does not establish one. What Section 12.1 states is the **competence** — what someone must be able to do. Whether that competence is ever formally certified, and by whom, is a question for a field with practitioners in it, not for a methodology at version 1.0.

13. The research agenda

The open questions of the discipline are maintained in the teaching primer, *Interpretability Engine Optimization: An Introduction*, Part 8. As of this version they include: measuring interpretability outside commerce and beyond the product object · separating representation effects from model effects across model families · standardizing gold labels and inter-rater reliability across instruments · measuring content quality at the structural layer · the economic chain end to end · interpretability debt over time · the ablation extension of the counterfactual test · and whether interpretation-for-task is a distinct layer or a refinement of machine-understandability.

Results that contradict the positions in this document are as welcome as results that support them.

14. Relation to adjacent practices

IEO does not claim that these techniques — or the underlying ideas of machine-actionability, semantic interoperability, task-specific fitness, provenance, or ontology evaluation — are new. It proposes to organize them around a specific operational objective: **whether a machine can correctly resolve a source representation for a declared decision context, before downstream retrieval, recommendation, comparison or action.** IEO's contribution is the framework by which that objective is specified, diagnosed, remediated, validated, measured and re-measured across implementation stacks and domains.

Close precedents exist for individual components and partial combinations of this objective. The claim of IEO is therefore methodological integration and scope, not priority over its constituent techniques or concepts.

Four bodies of work stand closest to the core of this methodology and are acknowledged as its principal antecedents:

Ontology engineering and evaluation. Competency questions — natural-language questions a representation must be able to answer — have been used as both design requirements and evaluation criteria since Grüninger and Fox (1995), and remain central to the field (Alharbi, Tamma, Payne & de Berardinis, 2026, *AI Magazine*, doi:10.1002/aaai.70054). Ontology evaluation has long addressed performance on the tasks a representation is designed for. The difference in scope: that literature evaluates whether *an ontology* adequately represents a domain; IEO's object may be any source representation — a page, a feed, a record, an API response, a document — and the discipline is stack-agnostic.

FAIR and machine-actionability. The FAIR principles (Wilkinson et al., 2016, *Scientific Data* 3:160018) define machine-actionability as an agent identifying what an object is, determining whether it is useful within the context of its current task, determining constraints on use, and taking appropriate action. The FAIR Maturity Indicators and Evaluator (Wilkinson et al., 2019, *Scientific Data* 6:174) supply objective, automatable assessment. Machine understanding, task context, use constraints and downstream action therefore do not originate with IEO.

Semantic interoperability. ISO/IEC 30182 and clinical terminology standards define the requirement that systems exchange information and correctly interpret its meaning, for which meaning must be agreed, consistent and clearly expressed. The objective of correct machine interpretation of meaning predates this methodology.

Contextual and semantic data quality. Data quality has been defined as fitness for use by data consumers since Wang and Strong (1996, *Journal of Management Information Systems* 12(4)), and later work has developed frameworks for semantic data quality based on intended use, incorporating provenance and pre-analysis assessment (Razzaghi, Greenberg & Bailey, 2022, *Learning Health Systems* 6(1)). The question "*interpretable for what?*" is therefore not new either.

A full positioning of IEO against each of these bodies of work — where they intersect and where each ends — is the subject of a separate paper.

A note on terminology. A taxonomy is not an ontology. A collection of records is not a knowledge graph. A database query is not retrieval-augmented generation. This methodology uses each term only where it applies, and practitioners are asked to do the same.

15. How this methodology evolves

An open discipline has to state how it changes, or it is one organization's document with an open licence attached.

Versioning. This methodology is versioned. Each version states what changed and why. Previous versions remain available exactly as published; corrections appear in a superseding version with an explicit correction record. Nothing is edited silently, so that anyone who cited an earlier version can see what moved.

Custodianship. NSOLVIA Research maintains this document as a public reference. Maintaining the reference does not restrict independent practice, research, implementation, or development of the discipline, and nothing in this methodology requires NSOLVIA's tools, participation or agreement to practice.

How to propose a change. Corrections, contradicting evidence, alternative formulations and new domain profiles are welcome. The mechanism is publication: publish the work, cite what it corrects, and it enters the record. A methodology that could only be changed by its author would not describe a discipline.

What would trigger a major revision. A replication that contradicts the counterfactual result. An independent instrument that produces incompatible measurements. A published domain profile that the framework cannot accommodate. Evidence that the economic chain does not hold. Any of these would require this document to change, and that is the point of stating them.

16. Limits

What IEO does not do.

It does not make a machine prefer an offer. It does not rank, surface or promote. It does not decide — it prepares interpretation for systems that decide. It does not generate information the source does not support. It does not replace the disciplines that operate downstream of interpretation, and it does not replace the domain expertise that supplies the truth it makes interpretable.

What IEO has not yet earned.

It is not an established field with an independent literature. Its terminology is not yet in general use. There are no industry standards for its instruments, and independent instruments have not yet been published. Its causal chain to commercial outcomes — that better interpretation produces measurably better business results — is stated as a measurable hypothesis and has not been demonstrated end to end. Evidence exists in one domain, concentrated in one of that domain's three objects.

Why these are stated here. A discipline matures through a known sequence: definition, methodology, metrics, benchmarks, causality, models of return. This document is at the second step. Declaring where it stands is not a concession; it is the practice applied to itself.

17. Terminology

The canonical definitions of every term used here are maintained in the public glossary at **nsolv-[a.com/ieo/glossary](https://nsolv.a.com/ieo/glossary)**. For this version of the methodology, **the methodology governs**; the glossary is kept aligned with it. A substantive change to a definition is published as a new version of the methodology, not as a silent edit to the glossary — so that anything citing this version means what it meant when cited.

Two framings are canonical and are quoted verbatim:

IEO optimizes understanding.

Semantic discovery is downstream of interpretation.

18. On openness

NSOLVIA proposes Interpretability Engine Optimization as an open discipline. This methodology is published under CC BY 4.0. It may be reproduced, adapted, taught, and built upon, including commercially, with attribution.

What is open: the discipline, this methodology, the commerce domain profile, the glossary, the measurement terms, the tests.

What is not open, and is not described here: how any particular organization — NSOLVIA included — implements the discipline. The methodology specifies the objective. The terrain is crossed by whoever practices.

Version history

Version 1.0 — 5 September 2026. First public methodology. DOI: 10.5281/zenodo.22416784. Companion to the founding paper (Version 1.2; concept DOI 10.5281/zenodo.22104280).

Interpretability Engine Optimization (IEO)

The UNDERSTAND Layer

