



Confidence by Design: Practical Patterns for Trustworthy AI Research Systems

Aide Solutions LLC
July 2026

This article was developed with the support of an AI-assisted writing system designed for scientific and professional communication materials.

Table of Contents

Table of Contents.....	2
Section 1: From Argument to Existence Proof.....	5
What the Series Established.....	5
The Demonstration Vehicle.....	5
What This Paper Does.....	6
Section 2: The Architecture of a Trustworthy Modeling System.....	6
Why Specialization Creates Checkability.....	7
Gates at Judgment Junctures.....	8
The Assurance-Artifact Principle.....	8
Section 3: Four Properties as Architectural Decisions.....	9
Transparency and Traceability: Legibility Built In.....	9
Selective Attention and Calibrated Human Oversight: Directing Judgment.....	10
The Interlocking System.....	11
Section 4: Independent Cross-Verification: Double Programming at the AI Level.....	11
The Verification Machinery.....	12
Convergence of Conclusion as the Acceptance Standard.....	12
Section 5: Evaluating the Built System.....	13
The Six Dimensions, Each Answerable from the System's Own Record.....	14
The Human-AI Modeling System as the Unit of Evaluation.....	15
What This Architecture Does Not Claim.....	15
Section 6: The Pattern Generalizes, and What It Means.....	16
What This Means for Three Audiences.....	17
The Arc Is Complete.....	17
References.....	17

Executive Summary

Health economic modeling sits at the intersection of scientific rigor and consequential decision-making, yet the tools used to evaluate AI-assisted modeling systems have not kept pace with the complexity those systems produce. Replicating a single ICER to within a narrow tolerance is a tractable benchmark, but it answers almost nothing about whether the structural choices, evidence selection, and analytical assumptions embedded in a model are defensible. This white paper closes a four-part series by moving from argument to demonstration, showing that a production system built on explicit trustworthiness principles is achievable today and describable in sufficient detail for professional evaluators to assess.

The series established three prior foundations. The Validation Gap identified the problem: as modeling outputs grow more complex and design-driven, traditional accuracy metrics capture a shrinking share of what determines trustworthiness. Beyond Replication proposed convergence of conclusion across four dimensions (treatment ranking, ICER range relative to willingness-to-pay threshold, primary uncertainty drivers, and CEAC probability) as the appropriate acceptance criterion, drawing on Mount Hood Challenge data in which ten independent modeling groups produced net monetary benefit variation of up to £23,696 while agreeing on treatment ranking. That paper also established the human-AI modeling system, encompassing the AI outputs, the analyst's judgment, the governing workflow, and the documentation making both auditable, as the correct unit of evaluation. The Confidence Layer translated evaluation philosophy into design philosophy, defining four properties that constitute a confidence layer when embedded in a system rather than applied after the fact: transparency, traceability, selective attention, and calibrated human oversight.

This paper demonstrates that those properties are buildable. The example is HEORACLE, a multi-agent system developed by Aide Solutions that takes a decision problem and published evidence base through four sequential modules: concept development, model specification, independent dual coding in Excel and R, and technical reporting. The architecture treats specialization as the basis of checkability. Because each agent is responsible for a bounded task, paired validation steps can be narrow, purposeful, and complete rather than exhaustive and approximate. Expert review gates are positioned at junctures where scientific judgment is genuinely required, not distributed uniformly across the pipeline. This placement reflects the principle of calibrated human oversight: the analyst drives scientific decisions while the system handles labor-intensive intermediate steps.

Each of the four confidence-layer properties is expressed as a load-bearing architectural decision. Transparency is implemented through live narration during generation, producing reasoning records as byproducts of the work rather than

retrospective summaries. Traceability is enforced through alignment logs linking every output element to the specification and evidence from which it derives. Selective attention is operationalized through uncertainty flagging that directs expert scrutiny where it is actually needed. Calibrated oversight is enforced structurally: the orchestration layer requires that each gate be satisfied before the pipeline advances.

For HEOR, Medical Affairs, and Market Access professionals who evaluate AI-assisted research, the practical implication is direct. Assurance in these systems depends on architecture, not assertion. The patterns described here are auditable, reproducible, and available for evaluation against the six dimensions developed in the prior paper. Organizations building or procuring AI-assisted modeling capabilities should demand this level of structural transparency as a baseline expectation, not an advanced feature.

Section 1: From Argument to Existence Proof

Three papers in this series built a case for rethinking how trustworthiness is constructed in AI-enabled health economic research. This paper closes the arc differently: by showing the philosophy built.

What the Series Established

The Validation Gap named the core problem [1]. As research outputs grow more complex and design-driven, traditional accuracy metrics capture a shrinking slice of what determines whether a model is trustworthy. Replicating a single ICER to within 1% of a published value [7] is a technically tractable benchmark; it answers almost nothing about whether the structural choices, evidence selection, and analytical assumptions embedded in that model are defensible. The gap between what can be measured and what matters grows with output complexity.

Beyond Replication replaced the inadequate standard with a better one [2]. *Convergence of conclusion* across four dimensions (treatment ranking; ICER range relative to willingness-to-pay threshold; primary uncertainty drivers identified by DSA and PSA; CEAC probability) is the appropriate acceptance criterion, because it reflects what decisions actually turn on. The Mount Hood Challenge illustrated why this matters: ten independent modeling groups working from identical inputs produced net monetary benefit variation of up to £23,696, yet agreed on treatment ranking [4]. Numerical identity across intermediate cells was neither achieved nor required. *Beyond Replication* also established that the *human-AI modeling system as the unit of evaluation* is the right frame: the AI's outputs, the human analyst's judgment, the governing workflow, and the documentation making both auditable must all be in scope.

The Confidence Layer translated evaluation philosophy into design philosophy [3]. Four properties, when embedded rather than bolted onto finished work, constitute a *confidence layer*: **Transparency** (reasoning and decision paths visible within the workflow, not reconstructed after the fact); **Traceability** (every output element linked to what informed it, chains followable without reconstruction); **Selective attention** (the system signals where it is on solid ground and where it is uncertain, directing expert scrutiny where it is actually needed); and **Calibrated human oversight** (the pipeline pauses at junctures where expert judgment is genuinely required, not uniformly throughout). The paper's central claim was that these properties are buildable today, and buildable embedded.

The Demonstration Vehicle

This paper shows the claim is true.

HEORACLE is a multi-agent system built by Aide Solutions. It takes a decision problem and a published evidence base through four sequential modules: concept development, model specification, model coding (independently in both Excel and R), and technical reporting. Expert review gates are positioned throughout, at the junctures where scientific judgment is required. The system was designed to embody the *confidence layer* architecture described in *The Confidence Layer*: each property is expressed as a load-bearing architectural decision,

verification runs alongside the pipeline rather than after it, and assurance artifacts are produced as byproducts of doing the work.

The table below maps each prior paper's contribution to its expression in the demonstrated system.

Series Paper	Contribution	Expression in this Paper
<i>The Validation Gap</i>	Named the problem	Motivates the architectural priorities
<i>Beyond Replication</i>	Established the evaluative standard	Convergence of conclusion as acceptance criterion; six audit dimensions as evaluation frame
<i>The Confidence Layer</i>	Defined the design philosophy	Four properties as load-bearing architectural decisions

The pattern described here stands independently of any one implementation. HEORACLE is the implementation that shows the pattern is real and producible in a production system for health economic modeling.

What This Paper Does

The remainder of this paper walks through the architecture (Section 2), examines each of the four properties as a specific architectural decision (Section 3), describes the cross-verification mechanism that delivers *convergence of conclusion* in practice (Section 4), applies *the six evaluation dimensions* from *Beyond Replication* to the built system (Section 5), and considers where the pattern generalizes and what it means for evaluators, builders, and the profession (Section 6). The argument in the prior papers was conceptual. The argument here is the demonstration.

Section 2: The Architecture of a Trustworthy Modeling System

Multi-agent architecture is often described as a capability story: more agents, more parallelism, more throughput. The design logic here runs differently. In HEORACLE, the multi-agent structure serves trustworthiness first. Specialization of agents enables specialization of validation, and the orchestration layer converts that principle into an enforced property of the system.

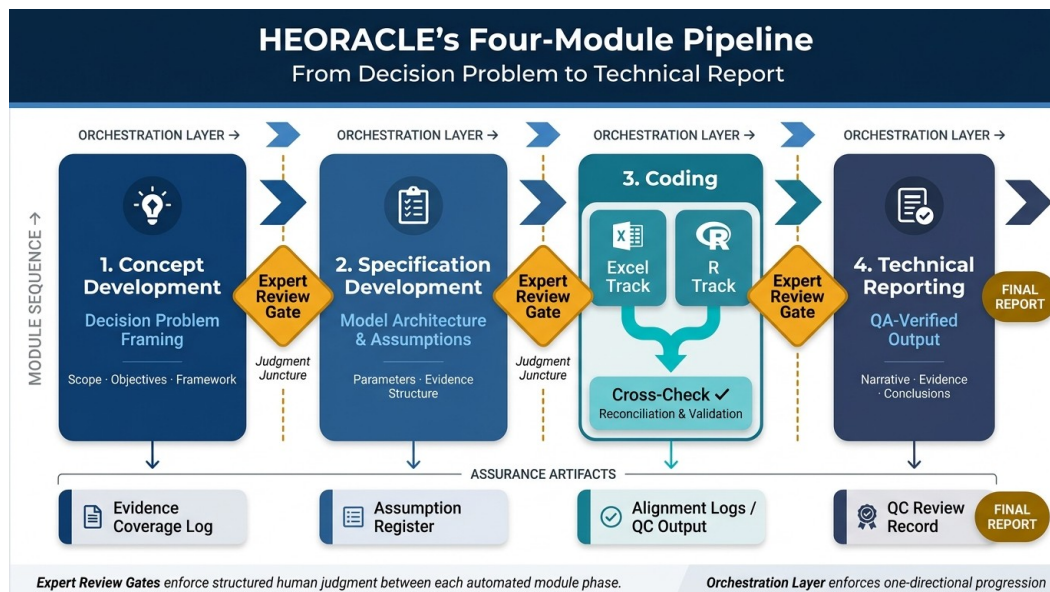


Figure 1. HEORACLE's Four-Module Pipeline: From Decision Problem to Technical Report

Why Specialization Creates Checkability

When a single model-building process produces a monolithic output, verification requires examining the whole artifact. Reviewers face an undifferentiated surface. When that same process is decomposed into discrete agents, each responsible for a bounded task, verification becomes tractable. Each agent's output has a defined scope; a paired validation step can be narrow, purposeful, and complete rather than exhaustive and approximate.

This is the architectural basis of *the confidence layer* in practice [3]. A concept-development agent responsible for scoping the decision problem produces an output that can be checked specifically against clinical relevance criteria and available evidence. A specification-development agent responsible for translating that concept into a formal model structure produces an output that can be checked specifically for structural soundness and evidence alignment. The coding agents, operating independently in Excel and R, produce outputs that can be checked specifically for computational fidelity to the specification. The reporting agent produces an output that can be checked specifically for alignment with the underlying model.

At each stage, the verification step is narrower than the generating step. That narrowness is deliberate. A narrow check is a completable check.

The orchestration layer enforces this architecture. Agents cannot skip gates. The progression from one module to the next requires that the verification checkpoint for the preceding module has been satisfied. This is an engineered constraint, not a procedural aspiration.

Figure 1 shows the four-module pipeline through which a decision problem and its evidence base are progressively structured, coded, and reported. Figure 2 shows the architecture in

full: agents and their paired validation counterparts, the gates at which expert decisions are required, and the points at which assurance artifacts are emitted.

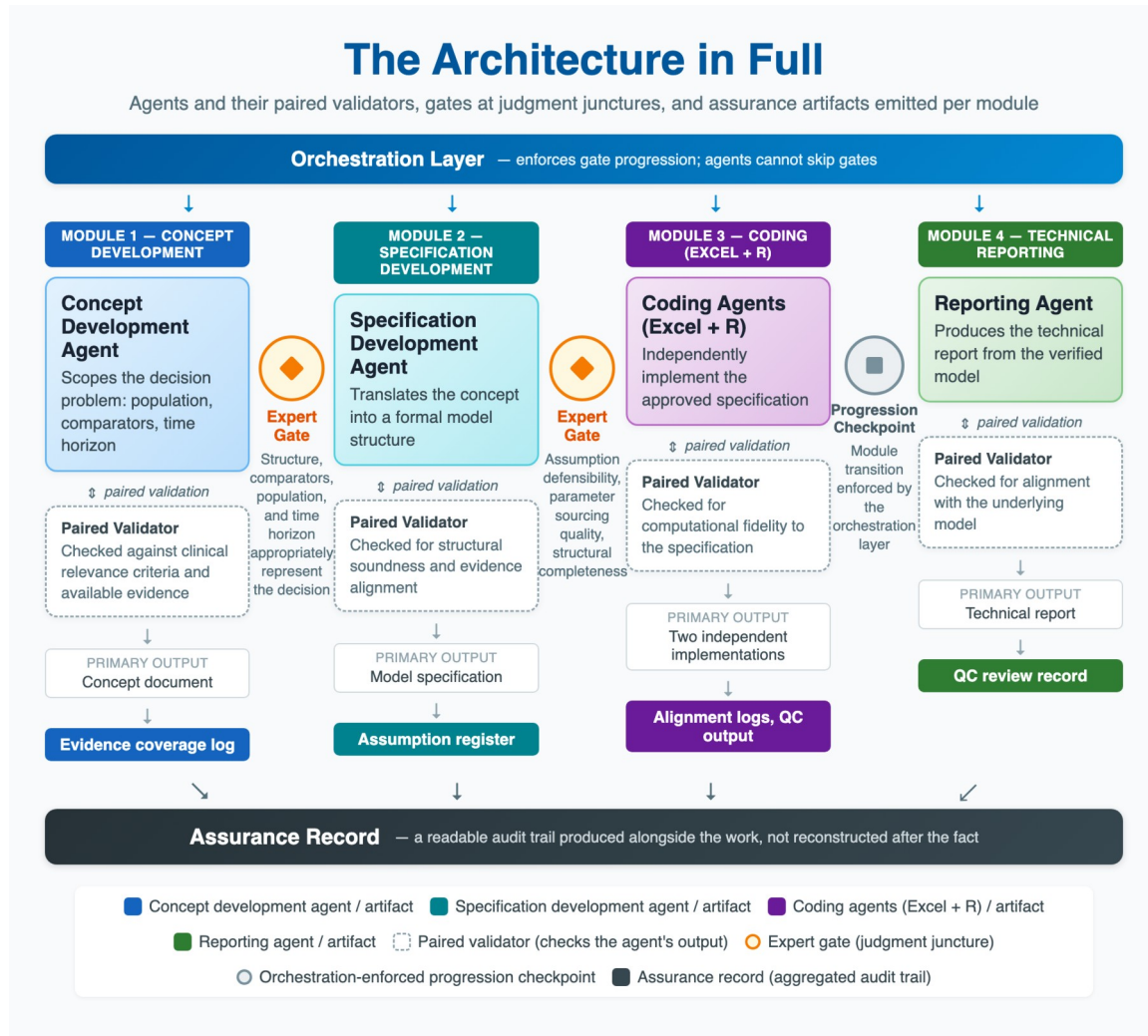


Figure 2. The architecture in full: production agents, paired validators, expert gates, and per-module assurance artifacts.

Gates at Judgment Junctures

The four modules, concept development, specification development, coding, and technical reporting, form a progressive structuring sequence. Each module takes the output of the previous module as its input and increases the level of formalization. A decision problem becomes a concept document; a concept document becomes a model specification; a model specification becomes two independent coded implementations; the implementations become a technical report with documented verification.

Expert gates are placed at the junctures in this sequence where scientific judgment is genuinely required. The gate between concept development and specification development is one such juncture: the expert determines whether the proposed model structure,

comparators, population definition, and time horizon appropriately represent the clinical and economic decision at hand. This is not a quality check that can be automated. It is a scientific decision. The gate captures it.

The gate between specification development and coding is a second juncture: the expert reviews the specification for assumption defensibility, parameter sourcing quality, and structural completeness before the implementation phase begins. Errors identified here are corrected in the specification, before they propagate into code.

Gates are not placed uniformly throughout the pipeline. Placing gates at every intermediate step would eliminate the efficiency rationale for AI assistance and substitute the cognitive burden of uniform review for the cognitive burden of manual construction. The placement reflects the principle of *Calibrated human oversight* [3]: the expert drives scientific decisions; the system handles labor-intensive steps in between.

The Assurance-Artifact Principle

The *embedded* vs. *bolt-on* distinction from *The Confidence Layer* [3] has a concrete expression in this architecture. At each module, the system emits its own assurance record as it runs. The concept development module produces an evidence coverage log. The specification development module produces an assumption register. The coding module produces alignment logs documenting the relationship between the specification and each implementation. The reporting module produces QC output.

These artifacts are not generated retrospectively. They are byproducts of doing the work. The table below summarizes the artifact emission pattern.

Module	Primary Output	Assurance Artifact
Concept development	Concept document	Evidence coverage log
Specification development	Model specification	Assumption register
Coding (Excel + R)	Two independent implementations	Alignment logs, QC output
Technical reporting	Technical report	QC review record

This structure means that when an evaluator asks to examine the basis for any element of the model, the record exists and is readable. The audit is not a reconstruction. The assurance record was produced alongside the work, at the moment the work was done, and it reflects what the system actually did rather than what was subsequently recalled or summarized. That is the practical meaning of *embedded* in a production system.

Section 3: Four Properties as Architectural Decisions

The Confidence Layer defined four properties, Transparency, Traceability, Selective Attention, and Calibrated Human Oversight, and argued they should be embedded rather than bolted onto finished work [3]. What follows shows each property as a specific, load-bearing architectural decision in a production system. The figures referenced here provide annotated illustrations of these decisions in operation.

Transparency and Traceability: Legibility Built In

Transparency, in the canonical definition from *The Confidence Layer*, means reasoning and decision paths are visible within the workflow, not reconstructed after the fact [3]. The architectural expression in HEORACLE is a live narration mechanism: as each agent builds, it produces a running account of what it is doing, why it is making specific choices, and where it is drawing from. Revisions are diffable. When the specification-development agent modifies a structural assumption, the prior state and the changed state are both visible at the moment of use. The record is a byproduct of the generation process itself, not a summary produced after the fact.

For practitioners familiar with justifying modeling choices to HTA reviewers, this distinction is substantive. A payer or HTA body typically asks not "what did you conclude?" but "how did you get there?" A system that narrates reasoning during generation answers that question from its own record. A system that produces a finished artifact and then documents its reasoning relies on reconstruction, which is both unreliable and unverifiable.

Traceability extends this principle from reasoning to evidence. Every element of the model specification links to its source evidence. When an assumption about treatment-specific utility is entered into the specification, the link to the source publication, the relevant table, and the extracted value are embedded at the point of specification. Cross-article provenance is explicit: where an assumption draws on multiple sources, each source is identified, and the chain from specification element to underlying evidence is followable without reconstruction.

This is what makes work defensible in HTA submission. Reporting standards for AI-enabled economic evaluations, including CHEERS-AI, identify explicit documentation of evidence sources and AI-assisted steps as expected practice [6]. Traceability as an architectural decision delivers that documentation as a structural property of how the specification is built, not as a reporting exercise performed after modeling is complete. Figure 3 shows the live build progress with stage narration and diffable revision history illustrating Transparency. Figure 4 shows the specification assessment and cross-article reference structure illustrating Traceability.

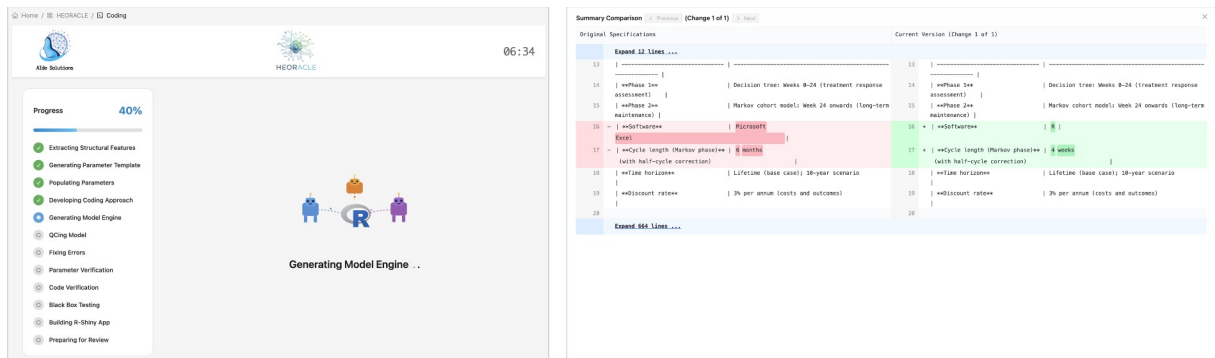


Figure 3. Live build progress with stage narration (left) and the diffable revision history view (right).

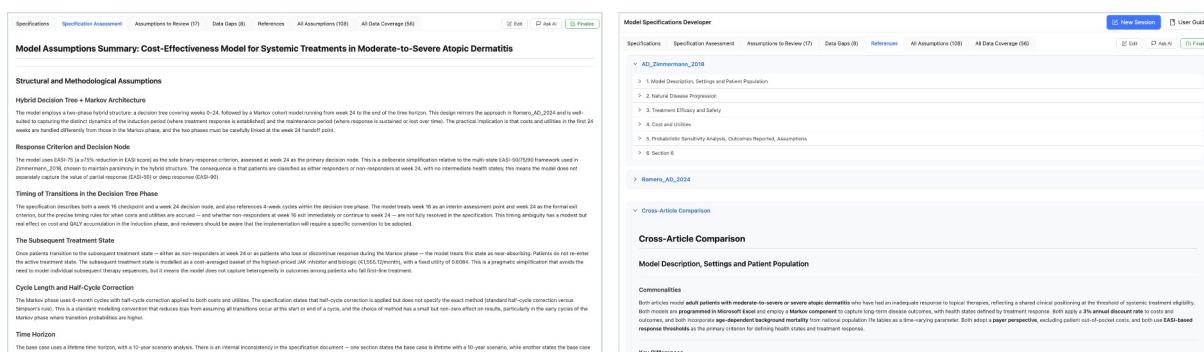


Figure 4. Specification assessment (left) and cross-article reference comparison (right).

Selective Attention and Calibrated Human Oversight: Directing Judgment

Selective attention addresses a practical problem that neither pure automation nor uniform review solves. Full automation eliminates the expert; uniform review eliminates the efficiency gain. Selective attention preserves both by having the system identify and surface what warrants scrutiny, rather than distributing review demand evenly across all outputs.

The architectural expression is an assumptions-to-review list generated during specification development. Assumptions carrying higher uncertainty, resting on single sources, or representing a departure from established modeling practice are flagged explicitly. Data gaps identified during the evidence-to-specification mapping are surfaced alongside the specification, annotated with the nature of the gap and its potential impact. The expert's attention is directed toward residual uncertainty, not distributed across what the system has already handled with confidence. An expert reviewing a 40-item assumptions register uniformly is less effective than an expert reviewing a targeted list of eight flagged assumptions with explicit annotations on why each warrants scrutiny.

Calibrated human oversight specifies where in the pipeline expert judgment is required and how it is captured. The pipeline pauses at judgment junctures, and the expert's decision is a structured input into a gate that governs whether the pipeline advances. The expert reviews the flagged assumptions list, considers the surfaced data gaps, and either approves the specification for coding or returns it with documented modifications. That decision and its basis become part of the assurance record. Figure 5 shows the assumptions flagged for review illustrating Selective attention. Figure 6 shows the gate structure at a judgment juncture, illustrating how the expert's decision is captured as structured input to Calibrated human oversight.

Model Specifications Developer

New Session

User Guide

Specifications

Specification Assessment

Assumptions to Review (17)

Data Gaps (8)

References

All Assumptions (108)

All Data Coverage (56)

Edit

Ask AI

Finalize

1

17 assumptions identified across 5 categories

Review each assumption, approve it, raise a question via the AI assistant, or override it with your own choice.

population_subgroups1 assumption · 1 pending

#1LOWPending Review

Approve

Question

Override

Alignment of trial population with modelled population

Efficacy data are sourced from clinical trials (JADE COMPARE, AD-Up, ECZTRA, BREEZE-AD) and a network meta-analysis (Silverberg et al.), with adjustments applied to align severe AD subgroup data with the modelled population; for comparators without direct severe-subgroup data, the moderate-to-severe NMA estimate is adjusted using the observed moderate-to-severe vs. severe ratio from JADE COMPARE.

RATIONALE

Direct head-to-head trial data across all seven comparators in a severe-only population are not available. Romero_AD_2024 addresses this by using JADE COMPARE data to derive a severity adjustment factor and applying it to NMA estimates for comparators without direct severe-subgroup data. This is a pragmatic approach to align trial populations with the modelled population.

ALTERNATIVES CONSIDERED

Alternative	Rationale	Trade-off
Use unadjusted NMA estimates (moderate-to-severe population) for all comparators	Avoids the assumption that the severity adjustment factor from JADE COMPARE is generalisable to all comparators; simpler and more transparent.	May underestimate response rates in the severe subgroup for all comparators; introduces a systematic bias if the modelled population is predominantly severe.
Restrict to comparators with direct severe-subgroup trial data	Eliminates the need for cross-trial severity adjustments; improves internal validity of efficacy inputs.	Substantially reduces the number of comparators that can be included; may exclude clinically important treatments (e.g., baricitinib, tralokinumab) from the analysis.
Use a separate NMA restricted to severe AD trials	A severity-specific NMA would provide the most internally consistent efficacy estimates for the severe subgroup.	May not be feasible due to limited number of trials reporting severe-only subgroup data; requires additional systematic review and NMA work.

IMPACT ON ANALYSIS

The severity adjustment methodology directly affects the relative efficacy of comparators without direct severe-subgroup data. If the JADE COMPARE-derived adjustment factor does not generalise to other treatments, the relative ranking of comparators could change. This is a key source of structural uncertainty and should be explored in sensitivity analyses.

SOURCES

- Romero_AD_2024 — Treatment Efficacy and Safety
- Compiled Model Specification — Treatment Efficacy and Safety

clinical_inputs_efficacy5 assumptions · 5 pending

Figure 5. Assumptions flagged for expert review, with rationale, alternatives considered, and approve, question, or override actions.

Health Economic Model Concept Development

New Session

User Guide

Evidence Coverage Review

Review the evidence gathered across research domains. Upload additional sources to fill gaps, then proceed.

Coverage by Domain

Category	Full Text	Abstract Only	OpenAlex	Web	HTA Domains	Coverage
Epidemiology & Demographics	4	0	0	4	0	Medium
Disease Progression & Outcomes	4	0	0	4	0	Medium
Quality of Life	8	0	0	7	1	High
Economic Burden	3	0	0	3	0	Medium
Treatment Landscape	10	0	0	10	0	High
Economic Evaluations	5	0	0	4	1	High

Supplement with Additional Sources (optional)

Supporting Publications

Enrich the evidence base with additional studies (epidemiology, burden of disease, utilities, cost data)

Drop files or browse

PDF, DOCX, DOC, MD, TXT

Intervention Document

The intervention that the model will focus on as the primary comparator

Drop file or browse

PDF, DOCX, DOC, MD, TXT

Continue with Current Evidence

Figure 6. An expert gate at a judgment juncture: the evidence coverage review requires an explicit decision before the pipeline proceeds.

12

The Interlocking System

The *Confidence Layer* argued that the four properties form an interlocking system rather than independent features [3]. The built system shows this rather than asserting it.

Property	What It Provides	Depends On
Traceability	Each flag links to its source evidence	Transparency (the chain was recorded, not reconstructed)
Selective attention	A targeted list of what warrants scrutiny	Traceability (each flag is only actionable if its evidence basis is followable)
Calibrated human oversight	A gate that acts on the flag list	Selective attention (the gate is only consequential if operating on a credible, evidence-linked list)

The flag list produced by Selective attention is only usable because each flag traces to evidence via Traceability. That evidence link is only credible because Transparency recorded the reasoning at the moment of generation. The gate governed by Calibrated human oversight is only consequential because it is acting on a substantive, evidence-grounded input. Remove any one property and the others degrade. Figure 7 shows the four properties as an interlocking dependency system.

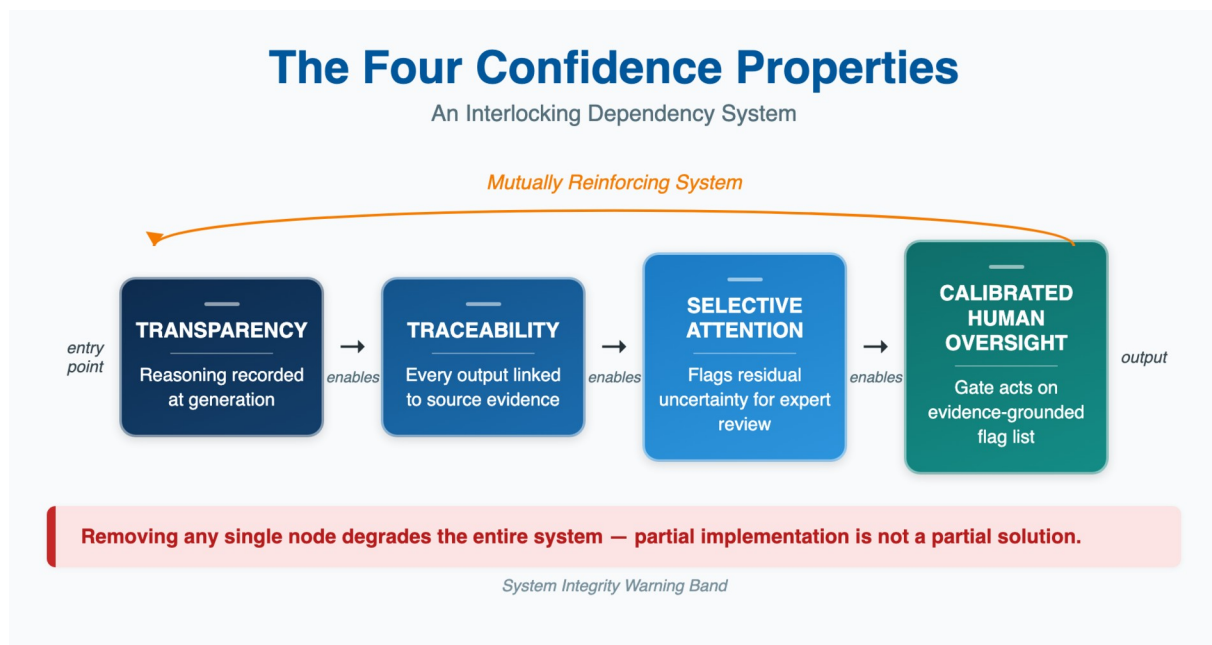


Figure 7. The four confidence properties as an interlocking dependency system.

For evaluators, the practical implication is that partial implementation is not a partial solution. An assumptions-to-review list without traceability is an unverifiable opinion. A

gate without a credible flag list is procedural theater. The architecture works because all four properties are present and mutually reinforcing.

Section 4: Independent Cross-Verification: Double Programming at the AI Level

Double programming is a well-established practice in health economic modeling. Two analysts build independent implementations from the same specification, then compare outputs. Discrepancies surface coding errors that neither implementation would expose on its own. The practice is recognized precisely because independent construction from a shared specification is a more reliable verification mechanism than reviewing a single implementation against itself.

HEORACLE applies this paradigm at the AI level. From a single approved model specification, two independent implementations are built: one in Excel, one in R. The implementations are constructed by separate coding processes without reference to each other's intermediate outputs. They are then cross-checked against each other and against the specification. Figure 8 shows the dual-output verification flow: from a shared specification, two independent implementations are built and cross-checked against each other and the specification.

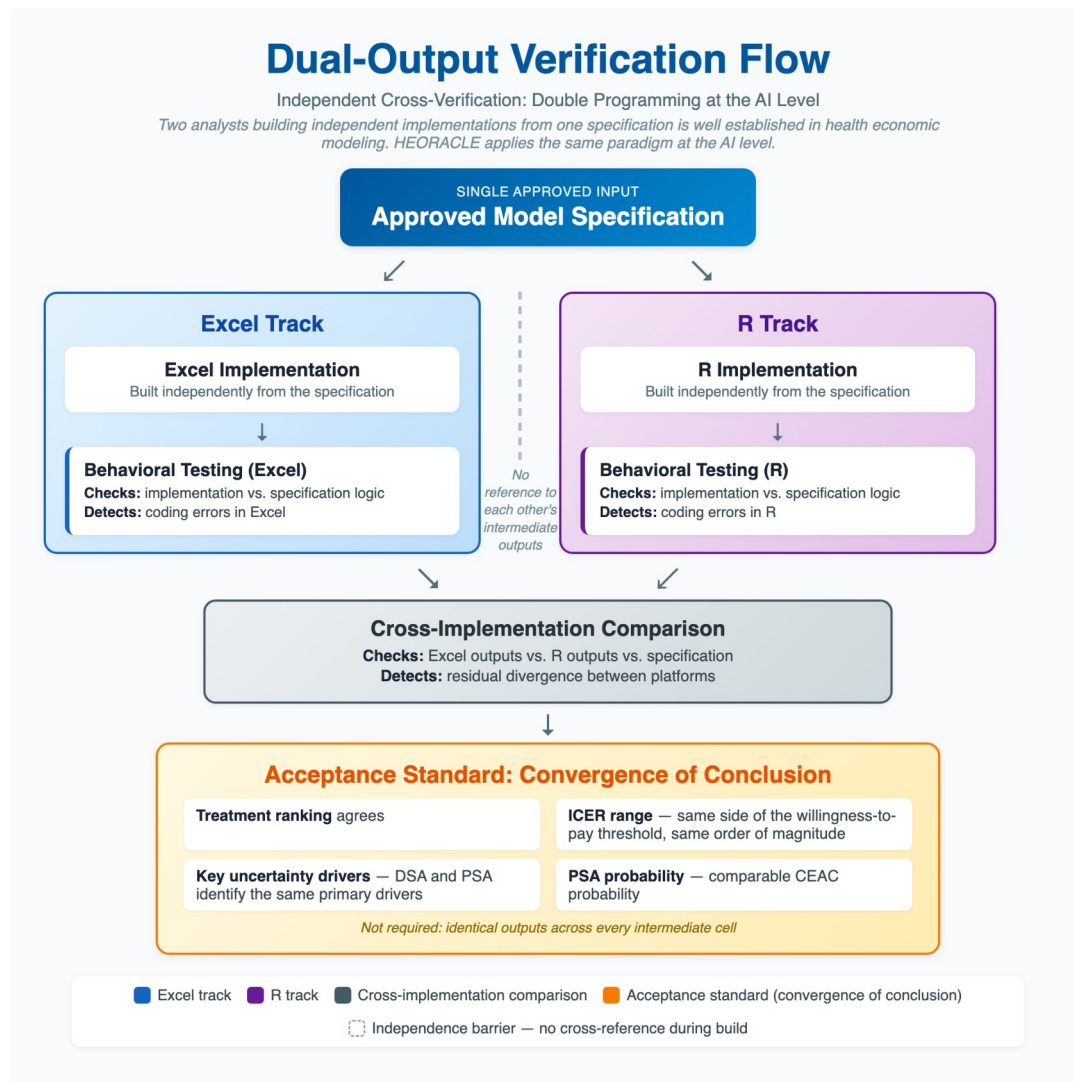


Figure 8. The dual-output verification flow: two independent implementations from one approved specification, behaviorally tested and cross-checked.

The Verification Machinery

Two mechanisms constitute the verification process operating within this architecture.

The first is behavioral testing. Before cross-implementation comparison, each coded model is tested against the specification directly. Outputs for defined input scenarios are compared against values derivable from the specification's structural logic. A model that behaves inconsistently with its own specification has a coding problem, regardless of whether its counterpart implementation agrees with it. Behavioral testing makes this class of error visible before the cross-check begins.

The second is the cross-implementation comparison itself. Once each implementation has passed behavioral testing, the Excel and R models are compared against each other and against the specification. Because the two implementations are constructed independently

from the same approved specification, their agreement carries evidential weight that neither implementation could generate about itself. Discrepancies are diagnosed to their source and documented; they are findings about the build, not signals to tune either implementation toward the other.

Alignment with published evidence is a different kind of question, and it is answered at the system level rather than inside any single build. Where the platform replicates published models, outputs are benchmarked against reported values in dedicated validation exercises of the kind documented in *Beyond Replication* [2], with deviations diagnosed to root cause and documented as validation findings. That evidence speaks to the accuracy of the system that produces the models. It is deliberately not a build-time loop that adjusts an implementation until it matches a published target: a model tuned to reproduce a number provides no independent evidence, and the verification logic of this section rests on independence.

The table below summarizes the verification sequence and what each step is designed to detect.

Verification Step	What It Checks	What It Detects
Behavioral testing (Excel)	Implementation vs. specification logic	Coding errors in Excel
Behavioral testing (R)	Implementation vs. specification logic	Coding errors in R
Cross-implementation check	Excel outputs vs. R outputs vs. specification	Residual divergence between platforms

Convergence of Conclusion as the Acceptance Standard

The acceptance criterion for cross-implementation agreement is *convergence of conclusion* across four dimensions, as established in *Beyond Replication* [2]: treatment ranking; ICER range (same side of the willingness-to-pay threshold, same order of magnitude); key uncertainty drivers (DSA and PSA identify the same primary drivers); and PSA probability (comparable CEAC probability). The implementations are required to agree on these dimensions. They are not required to produce identical outputs across every intermediate cell.

This distinction matters. The Mount Hood Diabetes Challenge Network 2022 is instructive on this point [4]. Ten independent modeling groups working from identical inputs produced net monetary benefit variation of up to £23,696, yet reached agreement on treatment ranking. That variation arose from legitimate structural and methodological differences among the groups, all working in good faith from the same evidence base. The groups disagreed numerically while agreeing scientifically.

The acceptance standard applied here reflects that same insight. Demanding numerical identity across every cell is *the replication trap*: it conflates technical verification with scientific evaluation, and it sets a standard that even expert human modelers do not

consistently achieve [2]. Convergence of conclusion targets what decisions actually turn on. If two independent implementations agree that treatment A dominates treatment B at a given willingness-to-pay threshold, that the same two parameters drive the uncertainty, and that the CEAC probability falls in the same range, the model has passed its verification test. Intermediate numerical differences are documented, not concealed, but they do not constitute failure.

This is what *Beyond Replication* means in practice [2]. The prior papers argued the standard; this architecture instantiates it. The cross-check is not designed to confirm that two implementations are identical. It is designed to confirm that they reach the same scientific conclusions, because conclusions are what HTA submissions, formulary decisions, and reimbursement negotiations depend on.

Section 5: Evaluating the Built System

An architecture that produces assurance artifacts as a byproduct of doing the work makes evaluation tractable in a specific way. The evaluator does not reconstruct what happened; the record was generated at the moment of generation. *The Validation Gap* framed the problem: as research outputs grow more complex and design-driven, traditional accuracy metrics capture a shrinking slice of what determines trustworthiness [1]. *Beyond Replication* established the evaluative frame to fill that gap: *the six evaluation dimensions*, grounded in AdViSHE's validation-assessment tool [5] for health-economic software, extended to cover AI-specific requirements [2]. What follows applies that frame to the built system. Figure 9 maps each evaluation dimension to the artifact that answers it.

The Six Evaluation Dimensions: Mapping Questions to Evidence

Evaluation Dimension	Evaluative Question	Assurance Artifact / Architectural Feature	Pipeline Module
1 Structural Soundness	Does the model structure appropriately represent the decision problem?	Concept document + model specification + expert gate records	Concept Specification
2 Assumption Defensibility	Is every material assumption justified, evidenced, and its uncertainty status disclosed?	Assumption register + documented expert overrides	Specification
3 Parameter Sourcing Quality	Are parameter values drawn from credible, traceable evidence?	Traceability links (specification element → source evidence); CHEERS-AI linkage	Specification
4 Sensitivity Coverage	Do DSA and PSA adequately cover structural and parametric uncertainty?	Built-in DSA/PSA + QC output from cross-implementation verification	Coding
5 Clinical & Economic Plausibility	Do outputs cohere with clinical reality and economic expectations?	Behavioral testing results + QC review record	Coding Reporting
6 Decision Relevance	Is the decision problem correctly framed — right comparators, right population?	Concept gate decision at pipeline entry	Concept
"The evaluator does not reconstruct what happened; the record was generated at the moment of generation."		Module legend: ● Concept ● Specification ● Coding ● Reporting	

Figure 9. The six evaluation dimensions mapped to the artifacts that answer them.

The Six Dimensions, Each Answerable from the System's Own Record

The table below maps each dimension to the artifact or architectural feature that answers it.

Evaluation Dimension	What Answers It
Structural soundness	Concept document + model specification + expert gate records
Assumption defensibility	Assumption register + documented expert overrides
Parameter sourcing quality	Traceability links from specification element to source evidence
Sensitivity coverage	Built-in DSA/PSA + QC output from verification
Clinical and economic plausibility	Behavioral testing results + QC review record
Decision relevance	Concept gate decision at pipeline entry

Structural soundness is answerable from two sequential records. The concept document captures the decision problem as scoped: population, comparators, time horizon, perspective. The model specification translates that scope into formal model structure. The expert gate between these stages records the reviewer's judgment that the structure appropriately represents the decision at hand. The gate record is a documented decision, not a timestamp.

Assumption defensibility is answered by the assumption register, which identifies every material assumption, its evidence basis, and its uncertainty status. Where the expert exercised an override, that override is documented with its rationale. An evaluator can read the register, locate the assumption, follow the traceability link, and determine whether any override was applied and why.

Parameter sourcing quality is answered by the traceability links embedded in the specification. Every parameter entry links to the source publication, the relevant table or figure, and the extracted value. CHEERS-AI identifies explicit documentation of evidence sources as a standard expectation for AI-enabled economic evaluations [6]. In this architecture, that documentation is a structural property of how the specification is built, not a post-hoc reporting exercise.

Sensitivity coverage is answered by the built-in sensitivity analysis and the QC output from the cross-implementation verification step. Deterministic and probabilistic sensitivity analyses are embedded in the model structure. The QC output documents the range of outputs examined and confirms that sensitivity runs executed as specified.

Clinical and economic plausibility is answered by the behavioral testing results and the QC review record. Behavioral tests check model outputs against the specification's structural logic and against values derivable from source publications. A model producing lower QALYs for a more effective treatment, or costs that violate structural assumptions, represents a surface-level failure behavioral testing is designed to catch.

Decision relevance is answered by the concept gate at pipeline entry. Before a model specification is developed, the expert determines whether the decision problem is correctly framed: whether comparators reflect the actual choice set, whether the population definition matches the intended decision context. A model can be structurally sound, well-parameterized, and fully verified while answering the wrong question. The concept gate is the check against that failure.

The Human-AI Modeling System as the Unit of Evaluation

The system being evaluated is *the human-AI modeling system as the unit of evaluation* [2]: the AI's analytical outputs, the expert's gate decisions and overrides, the workflow governing their interaction, and the documentation making both auditable. The expert's judgment exercised at each gate is a load-bearing component of the system, not external to it, and it is part of the system's record.

This matters for how dimensions are answered. Assumption defensibility is answered by the assumption register together with documented expert review of flagged assumptions. Parameter sourcing quality is answered by traceability links together with the expert's gate decision that the evidenced specification is fit to proceed. The architecture makes expert judgment visible and attributable, which is what allows it to be evaluated.

What This Architecture Does Not Claim

The assurance record makes error findable, attributable, and correctable. A documented expert override does not guarantee the override was correct. A traceability link does not guarantee the extracted value was appropriate. Passing behavioral tests does not guarantee the model structure is right for the decision problem.

The architecture provides the conditions under which trustworthiness can be assessed. The assessment remains the evaluator's responsibility. Claiming that architectural rigor substitutes for scientific judgment is accuracy theater. This architecture makes the record readable so that the evaluator's judgment, applied to that record, is well-informed rather than speculative.

Section 6: The Pattern Generalizes, and What It Means

Health economic modeling is a natural first domain for this architecture: outputs are complex, design-driven, and evaluated against demanding standards. But the pattern (specialized agents, paired validation, gates at judgment junctures, assurance artifacts as byproducts) is not specific to modeling. Any research task decomposable into discrete, checkable steps, requiring expert judgment at identifiable junctures, admits the same construction.

Evidence synthesis fits this description. Screening, extraction, quality appraisal, and synthesis are discrete steps, each with a defined scope and a checkable output. Expert judgment is genuinely required at specific junctures: inclusion decisions at the boundary of a PICO definition, quality appraisal of studies with mixed risk-of-bias signals, interpretation of heterogeneous effect estimates. A system built on the same four properties would narrate its screening logic as it runs (Transparency), link every extracted data point to its source (Traceability), surface borderline inclusion decisions for expert review (Selective attention), and pause at appraisal junctures requiring scientific judgment (Calibrated human oversight). The assurance artifacts would be a byproduct: screening logs, extraction records, decision documentation at each gate.

Evidence monitoring follows the same logic. Signal identification, triage, and relevance assessment decompose into discrete steps. Expert judgment is required when a signal crosses a threshold of potential significance. The gate exists; the architecture can enforce it and record what was decided.

Scientific writing is structurally similar. Drafting, citation verification, consistency checking, and editorial review are discrete steps. Expert judgment is required at specific points: whether a claim accurately reflects its underlying evidence, whether a conclusion

overreaches. A system with paired verification at each step and gates at judgment junctures produces a record of its own work alongside its output.

The generalization holds because the pattern is a research workflow pattern, not a modeling pattern. The four confidence properties are not specific to cost-effectiveness analysis; they are properties of how any complex, judgment-dependent research process can be made visible and verifiable.

What This Means for Three Audiences

For evaluators: Ask to see the assurance artifacts, not the demo. A demonstration shows what a system can produce under favorable conditions. Assurance artifacts show what the system produced during actual work, what was flagged, and what decisions were made at each gate. If a vendor cannot produce these artifacts, the confidence layer may not exist.

For builders: The four properties are buildable now, and embedded beats bolt-on. The confidence layer is an architectural decision, not a future capability contingent on more capable models. Bolt-on assurance, applied to a finished artifact, cannot recover the reasoning that generated it. Embedded assurance is recoverable by design.

For the profession: The field has spent considerable energy asking whether AI outputs can be trusted. The more productive question is whether the systems producing those outputs are designed for trustworthiness. Architecture answers that question in a way that output inspection alone cannot.

The Arc Is Complete

The Validation Gap named the problem: accuracy metrics capture a shrinking slice of trustworthiness as output complexity increases [1]. *Beyond Replication* established the evaluative standard: convergence of conclusion across four decision-relevant dimensions, audited through six evaluation dimensions, with the human-AI modeling system as the unit of evaluation [2]. *The Confidence Layer* defined the design philosophy: four properties, embedded rather than bolt-on, constituting a verifiable confidence layer [3]. This paper has shown the philosophy built.

The pattern is the proof.

References

- [1] Aide Solutions. The Validation Gap: What Accuracy, Sensitivity, and Specificity Cannot Tell You About AI Research Tools. April 2026. <https://aidesolutions.net/pub-the-validation-gap.html>
- [2] Aide Solutions. Beyond Replication: What It Actually Means to Evaluate an AI-Generated Health Economic Model. April 2026. <https://aidesolutions.net/pub-beyond-replication.html>
- [3] Aide Solutions. The Confidence Layer: A Design Philosophy for Trustworthy AI in Evidence Generation. June 2026. <https://aidesolutions.net/pub-the-confidence-layer.html>
- [4] Altunkaya J, Li X, Adler A, Feenstra T, Fridhammar A, Keng MJ, Lamotte M, McEwan P, Nilsson A, Palmer AJ, Quan J, Smolen H, Tran-Duy A, Valentine W, Willis M, Leal J, Clarke P. Examining the impact of structural uncertainty across 10 type 2 diabetes models: results from the 2022 Mount Hood Challenge. *Value Health*. 2024;27(10):1338-1347.
- [5] Vemer P, Corro Ramos I, van Voorn GAK, Al MJ, Feenstra TL. AdViSHE: a validation-assessment tool of health-economic models for decision makers and model users. *PharmacoEconomics*. 2016;34(4):349-361. <https://doi.org/10.1007/s40273-015-0327-2>
- [6] Elvidge J, Hawksworth C, Avsar TS, Zemplyeni A, Chalkidou A, Petrou S, Petyko Z, Srivastava D, Chandra G, Delaye J, Denniston A, Gomes M, Knies S, Nousios P, Siirtola P, Wang J, Dawoud D; CHEERS-AI Steering Group. Consolidated Health Economic Evaluation Reporting Standards for Interventions That Use Artificial Intelligence (CHEERS-AI). *Value Health*. 2024;27(9):1196-1205. <https://pmc.ncbi.nlm.nih.gov/articles/PMC11343728>
- [7] Reason T, Rawlinson W, Langham J, Gimblett A, Malcolm B, Klijn S. Artificial intelligence to automate health economic modelling: a case study to evaluate the potential application of large language models. *PharmacoEcon Open*. 2024;8(2):191-203. <https://doi.org/10.1007/s41669-024-00477-8>