

The Cost of Clarity

Scientific problem and applied-research method

Estimating the cost and risk of the minimal information an initiative needs — before commitment.

OÜ JUBAP · Tallinn, Estonia · July 2026
EIS Programme for Applied Research (RUP)

THEORETICAL ORIGIN · BACKGROUND, NOT PROJECT FOREGROUND

Where the hypothesis comes from

The constructs tested here were developed openly, before the project existed. They come from an international working group convened around AI integrity management and were published as field notes and working papers. That material is theoretical background. It is not project foreground.

PROGRAMME

Human Intelligence Debt

The economic framing. What an organisation pays when the information its decisions rely on was never established.

<https://tegrity.ai/series/human-intelligence-gap/>

RUP PROJECT

The Cost of Clarity

The field account behind the hypothesis. Which minimum information a transformation presupposes, and why nobody prices it.

<https://tegrity.ai/the-minimal-information-a-transformation-needs/>

STATUS OF THIS MATERIAL

- Openly published field notes and working papers. No peer review is claimed.
- Low technology readiness. Concepts and observations, not a validated method.
- No intellectual property and no commercial value are attributed to this material.
- Tegrity.AI is a publication channel. It is not an applicant or a project partner.**

The problem

Transformation loss is large, early commitments matter, and AI makes the same weakness sharper.

57%

of transformations miss value, timeline or both

BCG reports 57% of transformations failing to hit value, schedule, or both. A reported benchmark for that population, not a universal constant.

50%

of projects meet a modern success definition

PMI reports only half of projects meet a value-based success definition; 13% fail outright and 37% partially deliver.

75 / 15

cost committed early

NASA, using a DAU figure, states that during design about 15% of costs may be expended while about 75% of lifecycle cost is committed. This is a systems-engineering analogy, not an IT statistic.

>80%

reported AI failure estimate

RAND states “by some estimates” more than 80% of AI projects fail, about twice non-AI IT. RAND also identifies inadequate data as one of five leading root causes.

Each figure comes from a different population and definition, and none is a settled universal rate. Their evidential value is that they point in the same direction: transformation loss is large, it is committed early, and AI has not reduced it.

Why it matters: when an initiative is blocked for lack of owner, source, boundary or dependency clarity, the loss is not only its own budget. It drags dependent initiatives into waiting, escalation and rework across the roadmap.

And it is not improving

AI readiness and AI value-realisation sources disagree on definitions, but the direction is consistent.

13%

fully AI-ready “pacesetters”

Cisco’s third annual index reports a small pacesetter group at about 13% for three years; only 19% overall have fully centralised data.

17% → 42%

abandoning most AI initiatives

S&P Global / 451 Research reports companies abandoning most AI initiatives rose from 17% to 42%; average PoC scrapping reached 46%.

5%

measurable P&L impact

MIT NANDA reports only a small fraction of integrated GenAI pilots extracting measurable P&L value — a strict value-realisation measure, not a general failure rate.

AI does not create the problem; it removes the hiding place. A rationalisation may visibly stall. A model or agent keeps running and returns an answer — degraded, biased, or unverifiable — so the missing precondition is booked later as a data-quality defect.

Sources: Cisco AI Readiness Index 2025; S&P Global / 451 Research 2025; MIT Project NANDA 2025.

And yet the tools exist

Information and data readiness are not empty markets.

ACADEMIC / INSTITUTIONAL

- Data Readiness Levels
- ODI AI-ready data framework
- World Bank ODRA
- BADA / Vinnova

STANDARDS / MATURITY

- EDM Council DCAM
- DAMA-DMBOK
- Data-governance maturity
- ISO-aligned controls

CONSULTING / CLOUD

- McKinsey, Deloitte, Accenture
- PwC, KPMG
- Microsoft, AWS, Cisco
- Boutique services

SOFTWARE VENDORS

- Collibra, Alation, Atlan
- LeanIX/SAP, Ardoq, MEGA
- Celonis, Signavio, Apromore
- Informatica, UiPath

The scientific problem is not “why has nobody measured readiness?” It is why a mature and widely offered readiness category coexists with weak outcomes — and whether the issue is adoption at the commitment gate, instrument economics, attribution, or coverage.

The scientific problem

A persistent reported cause + a mature instrument category + weak outcomes.

Why?

If insufficient information readiness were simply a matter of not measuring it, a mature category of readiness assessment should have moved the outcome where it is applied before commitment.

The first part of the answer can be obtained without fieldwork: compare what ex-ante instruments specify with what ex-post accounts report as missing, contested, ambiguous or undeclared.

Sources: Lawrence 2017; RAND 2024; Suriadi et al. 2017; Ramasesh & Browning 2014.

What is the object we are measuring?

Not “data” in general — the minimal information a specific initiative presupposes.

Boundary

What is the object, population, process, application, case or portfolio slice?

Accountable owner

Who can decide, accept, fund, escalate and carry accountability?

Authoritative source

Which record, system, version or practice is canonical enough to act on?

Real dependencies

Which technical, process, supplier, regulatory and organisational links constrain movement?

“

Some of these facts are discovered. Others only exist once the organisation declares them.

Sources: Searle 1995/2006; Suriadi et al. 2017; process-mining and enterprise-architecture prior art.

Four candidate explanations

All four are plausible and non-exclusive. The project tests one first and keeps the others visible.

F1 Instrument economics

The assessment may cost more than the certainty it buys.

F2 Incentive structure

A reliable signal may be produced and still not change the decision.

F3 Attribution convenience

“The data was not ready” may be the easy story after the fact.

F4 Coverage

The tools measure available information, not the load-bearing missing set.

We do not pick a winner in advance. The design tests coverage (F4), measures instrument economics (F1), discriminates attribution convenience (F3) in the literature phase, and explicitly defers incentive structure (F2) as a larger longitudinal question.

F1 • Instrument economics

The assessment may cost more than the certainty it buys.

Mechanism: mass questionnaires and discovery work can duplicate what later acquisition must do anyway. If the signal adds little certainty per hour spent, teams rationally keep the assessment shallow.

Strong rival: even a better instrument may fail if unavailable information cannot be estimated from a sample at acceptable cost. This is the hard form of instrument economics (F1) and the mirror image of coverage (F4).



Measured here as acquisition cost per unit of uncertainty reduction.



Compared against random, stratified and centrality-only selection at equal cost.

No argument about “market irrationality” is needed. Instrument economics (F1) is an empirical output of the study.

F2 • Incentive structure

A good signal may be the hardest recommendation to act on.

A readiness signal is valuable only if it changes commitment behaviour. But acting on it may require a visible differentiated choice: rescope, resequence, stop, or run only the part where information exists.

That can be personally riskier than doing what comparable organisations do. This is an incentive equilibrium, not a claim that managers are careless or that governance is the cause.

This project does not claim governance is the cause.

Incentive structure (F2) is recognised and recorded where observable. Its market-wide contribution is deliberately deferred to a later study comparing high-performing and ordinary teams over time.

“ **Collected information can function as signal and symbol — not only as input to decisions.**

Sources: Feldman & March 1981; Meyer & Rowan 1977; DiMaggio & Powell 1983.

F3 • Attribution convenience

“The data was not ready” may be the safe post-mortem story.

Blame-neutral

Requires little forensic work

Fits almost any failure

Often partly true

If attribution convenience (F3) dominates, then after-the-fact accounts are weak ground truth, and a readiness instrument built from them would be aimed at the wrong target.

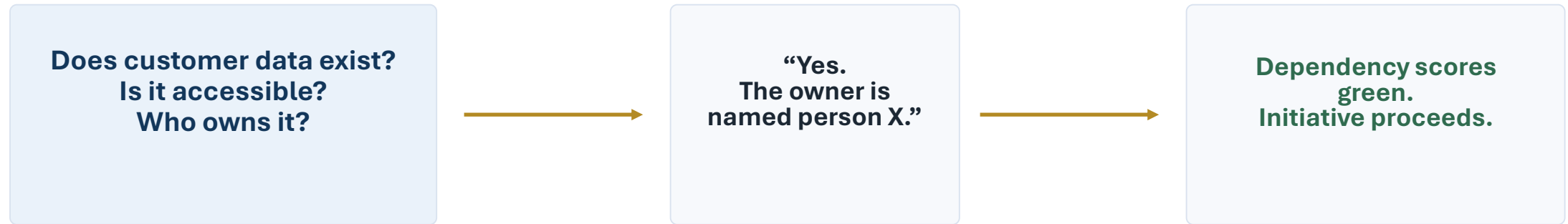
If attribution convenience (F3) is at work, the reports should look generic and low in detail, and much the same whatever the initiative was.
If a real coverage gap (F4) is at work, the reports should name specific missing items, and those items should differ by initiative type and be absent from the checklists.

The literature phase tests this before fieldwork, using preregistered bias criteria and blinded coding.

F4 • Coverage

The tool can answer its own question correctly — and still miss the load-bearing condition.

READINESS QUESTION



“ Nothing went wrong. The question was reasonable. The answer was truthful. The problem is what the question could not see.

Sources: Lawrence 2017; ODI 2025; EDM Council DCAM; Searle 1995/2006.

Coverage (F4) · What execution reveals

In prior field practice, a single customer-information object can attract dozens of overlapping ownership claims.

There is no master record. Several local processes maintain overlapping versions. Each local answer can be true. The multiplicity is the condition.

The initiative needs one canonical source. No fact exists yet about which source is canonical. It must be declared by an authority spanning the claimants.

This is not a retrieval problem. It is an organisational decision cost, before transformation begins, carrying real risk.

Why coverage (F4) is measurable: the multiplicity leaves traces — contradiction, authority dispersion, disagreement, non-response, retrieval friction and unclear evidence provenance.

Sources: Searle 1995/2006; Suriadi et al. 2017; Polanyi 1966 and Nonaka & Takeuchi 1995 on the conventional sense of tacit and explicit knowledge.

The distinction that makes coverage (F4) measurable

Hidden is not the same as not yet decided.

TACIT exists, but is not written down

A determinate fact exists. Someone knows it or practices it. It can be recovered, although it may be difficult.

UNDECLARED does not exist yet

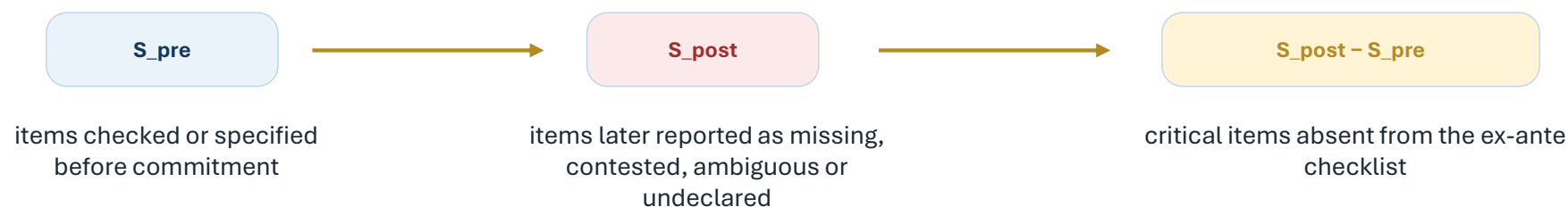
There is no value to retrieve or impute. The fact comes into being only when an authorised actor declares it.

Standard missing-data and imputation methods assume a true value exists. That assumption fails for undeclared information.

Sources: Searle 1995/2006; Rubin 1976 and Little & Rubin on missing-data assumptions; Polanyi 1966; Nonaka & Takeuchi 1995.

The measurable construct

The coverage gap is a set difference, not a mood.



The claim under test: the conditions reported missing after the event, but absent from the checklists used beforehand, form a substantial and structurally non-random set. If the set is empty, the coverage explanation (F4) is wrong. If it is generic and unpatterned, attribution convenience (F3) becomes stronger.

Sources: Suriadi et al. 2017; RAND 2024; ODI 2025; EDM Council DCAM; Lawrence 2017.

How the factors relate

They are not independent — and that makes a bounded study more informative.

Economics and coverage (F1 and F4) are two ends of one question. If the instruments sample what is available and miss what matters, then poor economics is the consequence of poor coverage. If instead no sample can estimate the missing set at acceptable cost, then the harder version of the economics explanation is right and no instrument can work.

Incentives and coverage (F2 and F4) are partly separated by the design. Comparing initiatives run by the same operator reduces incentive variation, but does not remove it. Incentive conditions are therefore recorded, not treated as solved.

Attribution and coverage (F3 and F4) can be told apart. Attribution convenience predicts generic, uniform reports. A coverage gap predicts specific items that differ by initiative type and are missing from the checklists. Both patterns are visible in literature that already exists.

Within-operator comparison: same organisation, similar governance, different initiative fates. This is a strategy to reduce confounding, not a magic control.

Correlation is a prediction, not a premise. The study must be able to find the factors independent — and report that.

Sources: Feldman & March 1981; Meyer & Rowan 1977; Searle 1995/2006; Ramasesh & Browning 2014.

Scope and hypotheses

The research question is deliberately narrow enough to be tested.

H1 • Coverage

The conditions reported missing after the event, but absent from the checklists used beforehand, form a substantial and specific set that differs by initiative type. Falsified if the reported missing set is already covered or too generic to code.

H2 • Transferability

Criteria from literature and high-performing teams can become quality-plan criteria another assessor can apply. Falsified if they remain tacit or non-auditable.

H3 • Prospective validity

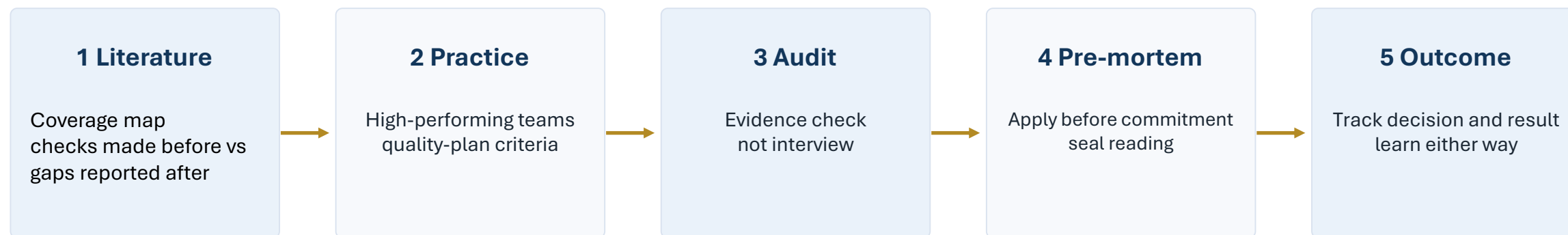
A sealed pre-commitment reading anticipates acquisition difficulty, rescope, delay, stop or under-delivery beyond basic complexity controls. Falsified if it adds no signal.

The incentive explanation (F2) is intentionally outside the primary scope. Testing it properly means comparing high-performing and ordinary teams across many organisations over time, which is a later and larger study.

Sources: Krippendorff 2018 on reliability; preregistration principles; technical annex (enclosed) for the full hypothesis set.

Method · from evidence to prediction

Five steps: read the literature, elicit practice that works, audit it against evidence, apply it before commitment, then observe what happens.



The starting capability exists in expert practice. The research result is a reproducible, transferable and economical instrument. Do not retrieve everything: estimate whether the critical information can be retrieved, elicited or declared at acceptable cost and risk. Nothing here requires a breakthrough; the contribution is configuration and validation.

Compare what is checked before with what is reported missing after.

Specificity	Generic phrase or item-level dependency / owner / boundary?
Patterning	Does it differ by initiative type, or appear everywhere?
Ex-ante visibility	Absent, implicit, explicit, unresolved or considered ready?
Information state	Retrievable, tacit/distributed, contradictory, undeclared, emergent?
Evidence strength	Contemporaneous record, triangulated account, perception or inference?

Bias control: criteria fixed before coding; coders blind to cross-tabulation; inter-coder reliability is reported. Generic attributions are not allowed to seed the battery.

Phase II · Practice that already works

Some teams already do this well. The question is whether their criteria transfer.

Select stakeholders leading teams with demonstrable, verifiable performance. The unit is the team and its programme record, not heroic self-description.

Ask concrete questions: on this named programme, what did you implement, decline, check, rescope or stop — and why?

Output = quality plan

steps · validation criteria · conformity · non-conformity · corrective action · preventive action

The conversion is the test. A criterion only its originator can recognise is not an instrument; it is tacit expert judgement.

Sources: ISO 10005:2018; ISO 19011:2026; Argyris & Schön on espoused theory and theory-in-use.

Phase II · Verification by audit, not interview

The account is treated as a claim. The claim is checked against evidence.

- 1 Stakeholder supplies a portfolio of prior initiatives.
- 2 A separate auditor selects which step is checked on which initiative.
- 3 Evidence required: record, artefact, decision, log or contemporaneous trace.
- 4 Unsupported criteria are corrected or removed.
- 5 Non-conformity is data — not participant failure.

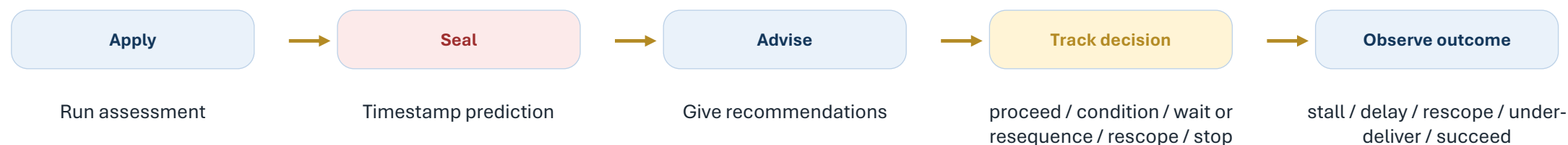
What this controls

war stories · self-selection ·
selective memory · originator bias ·
non-auditable tacit criteria

Judgemental audit sampling gives reasonable assurance about practice. It does not estimate population frequencies.

Phase III • Prospective pre-mortem

Apply the instrument before commitment, then observe what happens.



Primary result: agreement between the sealed reading and what later acquisition and outcomes show. This is different from claiming that acting on the reading caused improvement.

Decision field: record what was done after the reading. This makes incentive structure (F2) observable without turning it into the main claim.

Sources: Preregistration and prediction-sealing principles; technical annex (enclosed) for the prospective protocol.

Scientific controls

What makes it measurement, not opinion.

Two-source anti-circularity

Read gaps from independent sources where possible.

Reliability gate

No score is used before assessors agree.
Krippendorff $\alpha \geq 0.80$ as a hard gate.

Ground-truth calibration

Compare readings against acquisition evidence and realised outcomes.

Anonymity where needed

Anonymise the person, not the initiative. Preserve initiative traceability for validation.

Gate 0

If the instrument fails reliability or simple baselines, no substantive readiness claim is made.

Sources: Krippendorff 2018; ISO 19011:2026; preregistration principles.

How this can fail

Every branch is informative. That is the design, not a consolation.

No coverage gap	What is reported missing afterwards turns out to be already covered by the checklists, or the accounts are too generic to code. This would strengthen attribution convenience (F3) and correct our premise.
No transferable criteria	High performers have criteria, but they cannot be turned into quality plans another assessor can audit. This would measure directly where tacit expertise begins.
No reliable assessment	Assessors cannot agree on the same case. This is an instrument or execution failure, not a finding about readiness, and it is reported as such.
No incremental prediction	The readings add nothing beyond what initiative size and dependency count already explain. This would weaken the attribution of failure to information readiness.
Too costly	What is learned does not justify what it costs to learn it. This would support the harder version of the economics explanation: readiness may not be affordably knowable in advance.

“ A programme that cannot return a negative result measures nothing.

Sources: Preregistration and falsifiability principles; technical annex (enclosed) for the phase gates.

What would be known either way

The project buys a determinate answer — not a confirmation exercise.

Gap real + criteria transfer

Build the validated instrument and move to a larger validation/commercialisation stage.

Gap real + not detectable

Use staged commitment, reversible options and contracted discovery rather than pretending precision.

Gap not there

Rebuild the attribution literature: “information missing” was not the operative mechanism.

Tool works but no one acts

The incentive explanation (F2) becomes the next main study: how a reliable signal enters, or fails to enter, a commitment decision.

Business value: clearer proceed / condition / wait or resequence / rescope / stop decisions before the roadmap absorbs avoidable friction.

Sources: Technical annex (enclosed); staged-commitment and real-options literature on sequencing under uncertainty.

Scientific interfaces already discussed

The problem has been exposed to relevant specialists before being proposed as a study. Interest is not endorsement.

TalTech • expected R&D interfaces

Dirk Draheim • Ingrid Pappel • Riina Palu
Sijo Arakkal Peious • Vishal Choudhary

Tallinn University

Iren Irbe • tacit knowledge and TacitFlow
Katrin Niglas • mixed-methods and research design

University of Tartu

Anastasija Nikiforova • AI/data governance
Maaja Vadi • organisational behaviour
Fredrik Milani • BPM / process management

External validators / challengers

Rick Kazman • Wil van der Aalst • Josep Carmona
Javier Bécares • José Antonio Posadas García

Implementation / risk-market dialogue

Cybernetica AS: Dan Bogdanov • Aiko Adamson
Risk-management and trustworthy-data use cases

Interest is not endorsement. No participation, work package or cost is committed by any person or institution named here, and roles remain subject to Enterprise Estonia approval and subsequent agreements.

Source basis: public researcher profiles and July 2026 project correspondence available on request.

IP and franchise layers

The project publishes enough science to be testable, while protecting the operating system that makes it deliverable.



Commercial logic: the service can be delivered by experts and authorised partners before the central agentic module is complete. The franchise unit is the configured quality plan plus trained assessor, QA and protected calibration — not autonomous SaaS.

Sources: ISO 10005:2018; ISO 19011:2026; ISO/IEC 42001:2023; NIST AI RMF 1.0; EU AI Act 2024/1689; technical annex (enclosed).

Possible software modules

Software follows the validated criteria; it does not define them in advance.

1 · Criterion library

stores criteria, contexts, evidence rules and corrective actions

2 · Configuration engine

selects quality-plan scope by company and initiative type

3 · Evidence intake

collects, minimises, normalises and links documents, responses and traces

4 · Question selector

chooses next evidence request under cost, uncertainty and decision leverage

5 · Conflict detector

flags contradiction, authority dispersion and undeclared conditions

6 · Inference workspace

estimates burden, confidence, exposed value and readiness position

7 · Assessor workbench

keeps human approval, provenance, QA and audit trail

8 · Outcome feedback

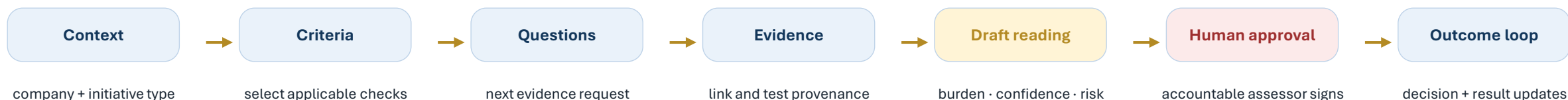
tracks decisions and later acquisition / delivery outcomes for calibration

Near term: expert-operated modules can be sold as controlled pilots. Next layer: central selection, inference and QA reduce cost and enable partner delivery.

Sources: Technical annex (enclosed); ISO 10005:2018 quality-plan structure; ISO 19011:2026 audit controls; NIST AI RMF; EU AI Act human oversight principle.

Agentic architecture — bounded orchestration

The proposed agentic layer helps select, assemble and trace work. It does not decide for the organisation.



What the agent may do

assemble plan · rank criteria · draft evidence requests · detect contradiction
· maintain traceability · suggest next observation

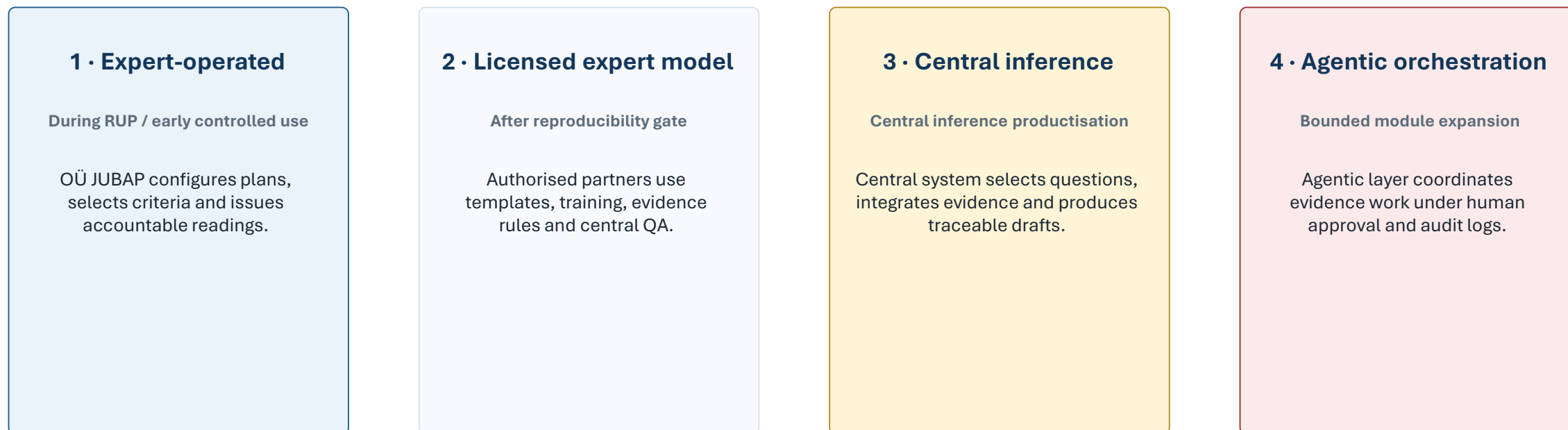
What it may not do

make organisational declarations · sign final conclusions · overrule assessor
judgement · act as unsupervised self-service decision authority

Sources: ISO/IEC 42001:2023; NIST AI RMF 1.0; EU AI Act 2024/1689 on risk management and human oversight; technical annex (enclosed).

Commercialisation path

Once the relevant module passes its validation gate, expert delivery can start before full automation; automation increases throughput.



Core message for EIS: the RUP validates a human-executable scientific instrument and structured assets for scale. Bounded assisted-processing modules may be included where justified; the production-grade multi-partner agentic system remains a productisation path, not a prerequisite for commercial service.

Sources: Technical annex (enclosed); ISO 10005/19011; ISO/IEC 42001 and NIST AI RMF for controlled AI-assisted operation.

Selected references · phenomena and current tools

Full URLs are included to support technical review.

BCG (2021). What Works—and What Doesn’t—in Transformation. <https://www.bcg.com/publications/2021/how-companies-implement-successful-transformation>

PMI (2025). Strategy-execution gap / project success research. <https://www.pmi.org/about/press-media/2025/new-pmi-research-reveals-strategy-execution-gap-is-undermining-transformation-and-how-to-close-it>

NASA SEH 2.5. Cost Effectiveness Considerations. <https://www.nasa.gov/reference/2-5-cost-effectiveness-considerations/>

RAND (2024). The Root Causes of Failure for Artificial Intelligence Projects. https://www.rand.org/pubs/research_reports/RRA2680-1.html

Cisco (2025). AI Readiness Index / Pacesetters. <https://investor.cisco.com/news/news-details/2025/Cisco-AI-Research-The-Most-AI-ready-Companies-Outpace-Peers-in-the-Race-to-Value/>

S&P Global / 451 Research (2025). Generative AI shows rapid growth but yields mixed results. <https://www.spglobal.com/market-intelligence/en/news-insights/research/2025/10/generative-ai-shows-rapid-growth-but-yields-mixed-results>

MIT Project NANDA / MIT Media Lab (2025). The GenAI Divide. <https://www.media.mit.edu/groups/nanda/overview/>

Selected references · instruments and prior art

The project is positioned against existing readiness categories, not against straw men.

Lawrence, N. D. (2017). Data Readiness Levels. <https://arxiv.org/abs/1705.02245>

Open Data Institute (2025). A framework for AI-ready data. <https://theodi.org/insights/reports/a-framework-for-ai-ready-data/>

World Bank. Open Data Readiness Assessment. <https://www.worldbank.org/en/data/opendatatoolkit/odra>

EDM Council. DCAM / Global Data Management Benchmark. <https://edmcouncil.org/innovation/research/benchmarks/>

Suriadi et al. (2017). Event log imperfection patterns for process mining. <https://doi.org/10.1016/j.is.2016.07.011>

van der Aalst et al. (2012). Process Mining Manifesto. https://doi.org/10.1007/978-3-642-28108-2_19

BADA / Vinnova (2017). Data Readiness for BADA. <https://ri.diva-portal.org/smash/record.jsf?pid=diva2:1170395>

Selected references · factors and method

For each factor and control logic.

Searle (1995). The Construction of Social Reality. <https://search.worldcat.org/title/31411549>

Searle (2006). Social Ontology: Some Basic Principles. <https://doi.org/10.1177/1463499606061731>

Feldman & March (1981). Information in Organizations as Signal and Symbol. <https://doi.org/10.2307/2392467>

Meyer & Rowan (1977). Institutionalized Organizations. <https://doi.org/10.1086/226550>

DiMaggio & Powell (1983). The Iron Cage Revisited. <https://doi.org/10.2307/2095101>

Fischhoff (1975). Hindsight ≠ foresight. <https://doi.org/10.1037/0096-1523.1.3.288>; Roese & Vohs (2012). <https://doi.org/10.1177/1745691612454303>

Ramasesh & Browning (2014). Knowable unknown unknowns. <https://doi.org/10.1016/j.jom.2014.03.003>

ISO 19011:2026. Guidelines for auditing management systems. <https://www.iso.org/standard/19011>

ISO 10005:2018. Quality management — Guidelines for quality plans. <https://www.iso.org/standard/70398.html>

Selected references · implementation and governance

For quality plans, auditability, AI governance and controlled agentic operation.

ISO 10005:2018. Quality management — Guidelines for quality plans. <https://www.iso.org/standard/70398.html>

ISO 19011:2026. Guidelines for auditing management systems. <https://www.iso.org/standard/19011>

ISO 9001:2015. Quality management systems — Requirements. <https://www.iso.org/standard/62085.html>

ISO/IEC 42001:2023. Artificial intelligence management system. <https://www.iso.org/standard/42001>

NIST (2023). Artificial Intelligence Risk Management Framework 1.0. <https://doi.org/10.6028/NIST.AI.100-1>

Regulation (EU) 2024/1689 — Artificial Intelligence Act. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>

OÜ JUBAP (2026). Technical annex, stakeholder map and EIS clarification reply of 24 July 2026 — project records available on request.

OÜ JUBAP · TALLINN, ESTONIA

Thank you

JubAp.EU

Strategic Intelligence & Transformation

Address **Akadeemia tee 42-11, Tallinn, Estonia**

E-mail **info@jubap.eu**

Telephone **+372 60 284 01**

Web **jubap.eu**



Scan for the project page