

From SECI to EASCI: Operationalising Tacit Knowledge Externalisation

Iren Irbe

Tallinn University, Estonia

irbe@tlu.ee

Abstract: Organisations rely on tacit knowledge in complex, high-stress work, yet externalising it is difficult. The SECI model identifies externalisation as the bridge between personal expertise and organisational memory, but treats it as a single phase and offers limited guidance when articulation is partial, cognitively constrained, or obscured by GenAI. This paper proposes EASCI (Experience-Articulation-Structuring-Consolidation-Innovation), which decomposes externalisation into three operations. Articulation reveals partial reasoning through anchored prompts; Structuring organises fragments into provisional knowledge artefacts through abduction with mandatory counterhypotheses; Consolidation subjects artefacts to peer review before organisational use. Experience and Innovation anchor the lifecycle, seven design principles bound GenAI to scaffolding roles, and a W3C PROV-O layer preserves human-AI attribution. A scoping review, qualitative interviews in high-stress public-sector contexts, and participatory design inform the framework. An investigative design case demonstrates feasibility under the EU AI Act and Law Enforcement Directive. Six falsifiable propositions and four study designs address a persistent gap: knowledge-management frameworks are widely adopted but rarely tested.

Keywords: Tacit knowledge externalisation, SECI, Generative AI, Knowledge management, Abductive reasoning, Design science research

1. Introduction

Michael Polanyi observed that “*we can know more than we can tell*” (Polanyi, 1966). The gap between what experts know and what they can explain has resisted management efforts to close it. Practitioners develop judgment, pattern recognition, and situational understanding through experience, yet much of that expertise remains resistant to documentation. When they leave, organisations lose not only procedures but also the tacit reasoning that made those procedures effective under ambiguity and time pressure. The consequences are acute in high-stress domains such as fraud investigation, emergency response, and legal practice, where decisions have legal or safety implications and must withstand later scrutiny (Collins, 2010).

Three decades of knowledge management research have not solved this issue. SECI (Nonaka and Takeuchi, 1995) remains the dominant account of tacit-explicit conversion, but it treats externalisation as a single stage. Nonaka and Konno’s *Ba* (1998) recognises the need for supportive conditions yet offers little guidance when articulation is partial, cognitively constrained, or distorted by automated expertise (Gourlay, 2006; Tsoukas, 2005; Anderson, 1982).

GenAI intensifies the externalisation problem. It can draft, summarise, and structure knowledge artefacts faster than organisations can validate them (Dwivedi et al., 2023), creating four risks: persuasive opacity, provenance erasure, validation bypass, and drift amplification. In regulated domains, the EU AI Act (2024) and *Law Enforcement Directive* (2016) therefore require AI-assisted decisions to remain traceable, contestable, and attributable to human oversight.

This paper introduces *EASCI* (*Experience-Articulation-Structuring-Consolidation-Innovation*), a framework that operationalises the SECI model by decomposing externalisation into three processes. *Articulation* reveals partial reasoning through structured prompts, *Structuring* organises fragments into provisional explanations through abduction with at least one counterhypothesis, and *Consolidation* subjects the result to peer validation before organisational use. *Experience* captures situated context, and *Innovation* governs versioning, review, and retirement. EASCI therefore does not replace SECI; it operationalises its externalisation phase.

The paper contributes in four ways:

1. the EASCI process model, which breaks down externalisation into five stages with clear quality gates;
2. seven design principles outlining how GenAI functions as a bounded elicitation scaffold rather than an autonomous knowledge creator;
3. empirical evidence from a scoping review of 55 sources, qualitative interviews (N=10) with public-sector practitioners, and participatory design (Irbe, 2025); and
4. an evaluation framework comprising six falsifiable propositions with clear rejection thresholds.

2. Background: Tacit Knowledge and SECI

Tacit knowledge includes skills, intuitions, mental models, and context-dependent judgments that experts use fluently but find difficult to explain (Hinds et al., 2001). Cognitive research explains why: expertise develops through repeated practice, as complex reasoning sequences are chunked into automatic patterns that operate below conscious awareness (Anderson, 1982). *Collins's taxonomy* (Collins, 2010) replaces Polanyi's single concept with three distinct types: *Relational Tacit Knowledge (RTK)*, which could theoretically be made explicit but remains tacit due to social barriers; *Somatic Tacit Knowledge (STK)*, embodied in motor skills; and *Collective Tacit Knowledge (CTK)*, held within social practices that no individual fully possesses.

SECI's four modes are not functionally equivalent. Only externalisation requires the challenging cognitive effort of converting implicit reasoning into communicable form, yet SECI treats all four modes as conceptually similar (Gourlay, 2006). Three limitations prompt EASCI. First, SECI's single "Externalisation" phase conflates three cognitively distinct operations. Second, Ba assumes supportive conditions arise naturally, which contradicts the scoping review's finding that limited socialisation opportunities (38.2% of studies) and trust deficits (25.5%) are among the most common barriers (Irbe and Ogunyemi, 2025). Third, SECI offers no mechanism for temporal governance: the spiral metaphor implies renewal without specifying how knowledge should be versioned, reviewed, or retired.

Epistemologically, EASCI integrates three traditions: a pragmatist account of inquiry (Dewey, 1938; Peirce, 1997), in which knowledge is validated through its consequences in practice; a constructivist account of social validation (Weick, 1995; Nonaka and Takeuchi, 1995), in which knowledge becomes organisational through shared sense-making; and a cognitive science account of expertise (Anderson, 1982; Sweller, 1988), which explains why tacit knowledge resists direct articulation. The framework does not privilege one tradition: pragmatism grounds the Articulation and Structuring stages, constructivism grounds Consolidation, and cognitive science grounds the stage-separation principle and the design of prompts.

Recent work has begun adapting SECI to AI-mediated knowledge work. One approach maps AI capabilities onto each SECI mode but remains descriptive (Zhang et al., 2025). The *GRAI* framework (Böhm and Durst, 2026) adds a machine-agent layer but lacks operational detail. A systematic review of 97 GenAI-KM studies found externalisation to be the SECI mode most often supported by GenAI tools (Costa and Fantini, 2025). Another approach extends Ba to GenAI through "Reflective Ba" as an AI-mediated cognitive environment (He et al., 2026). EASCI differs by specifying stage transitions, typed artefacts, gate criteria, and bounded AI roles.

3. Research Motivation

EASCI originated from design science research (Hevner et al., 2004), integrating three empirical inputs. A scoping review following the Arksey-O'Malley framework (Arksey and O'Malley, 2005) analysed 55 sources (2019-2024) on tacit knowledge transfer (Irbe and Ogunyemi, 2025). The most frequently reported barriers were limited socialisation opportunities (38.2%), lack of trust and willingness to share (25.5%), and reliance on experienced employees (23.6%). Trust has a dual role (a primary barrier when absent, a primary enabler when present) (Mayer et al., 1995) and motivates EASCI's consolidation stage.

Semi-structured interviews with public-sector practitioners (N=10) in investigative and analytical roles confirmed these findings and revealed two patterns. First, the what-vs-why gap: practitioners could more easily explain what they did than why; scenario-based and case-specific prompts uncovered more decision-relevant content, aligning with dual-process accounts of cognition (Eraut, 2004; Ericsson and Simon, 1993). Second, practitioners showed reluctance to share reasoning perceived as incomplete, consistent with causal ambiguity (the difficulty of identifying exactly why a practice works) as a barrier to knowledge transfer (Szulanski, 1996). An illustrative design case in investigative work demonstrated feasibility across all five stages, operating on air-gapped infrastructure in compliance with the Law Enforcement Directive.

4. The EASCI Framework

EASCI organises tacit knowledge externalisation as a five-stage lifecycle. Each stage produces a distinct artefact and is governed by explicit transition criteria that can be implemented as validation rubrics specifying what evidence, context, provenance, dissent, or lifecycle trigger must be present before a knowledge artefact can advance, return for repair, or be retired. In high-stress environments, where time pressure and cognitive overload shorten the reflection window, this staged structure is necessary rather than optional.

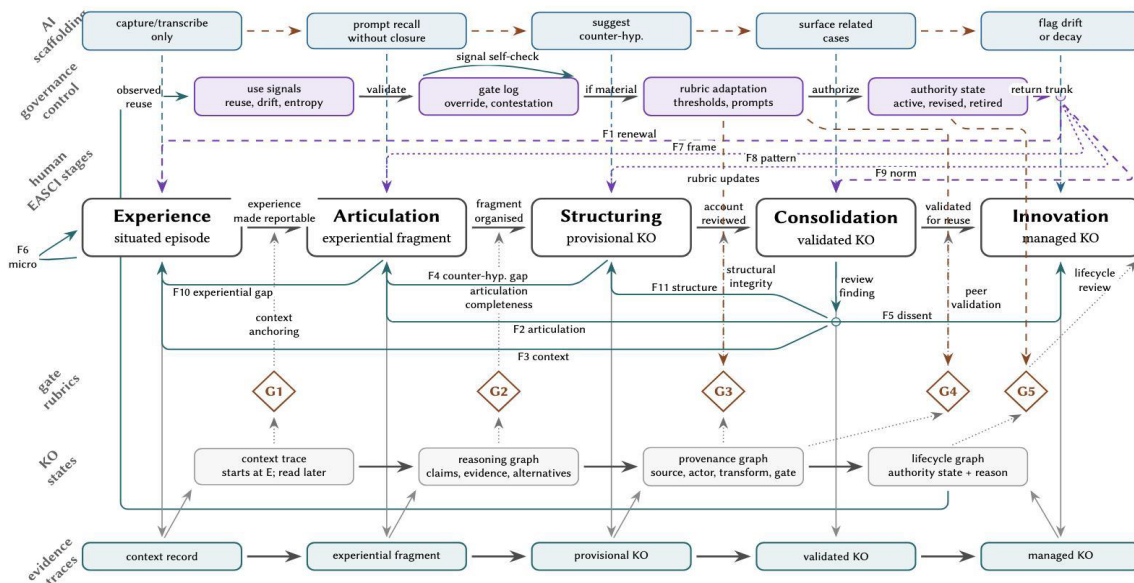


Figure 1: Integrated EASCI architecture

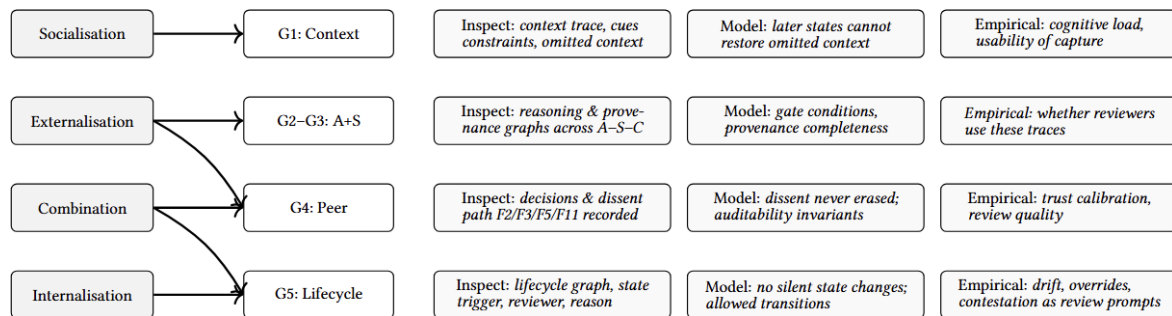


Figure 2: EASCI quality gates as a bridge between SECI and inspectable governance

Figure 1 shows the stable architecture: bounded AI scaffolding, a governance-control band, the five human EASCI stage states, gate predicates operationalised as validation rubrics, and accumulated context, reasoning, provenance, and lifecycle traces. Gates sit below stage transitions so accumulated traces can be inspected before a knowledge artefact advances. The governance-control band receives signals of reuse, drift, entropy, override, and contestation, determines whether they are material, and may update gate rubrics or lifecycle authority. These signals prompt review rather than serving as automatic evidence of benefit, harm, convergence, or quality improvement. The figure also distinguishes formal feedback-path families from local gate-failure actions, so repair, context recovery, and re-elicitation are not mistaken for additional feedback paths. Figure 2 then maps SECI to quality gates and inspectable governance. Four trace structures make both figures inspectable: the context trace, reasoning graph, provenance graph, and lifecycle graph.

4.1 Stage Descriptions

Experience captures conditions at the moment of action, before reflection or abstraction. Because expert knowledge is context-bound (Barsalou, 2020), decontextualised externalisation produces incomplete artefacts. Practitioners therefore record triggers, constraints, signals, and initial interpretations at the time of action, while AI is limited to prompting and metadata capture.

Articulation then reveals partial observations, hunches, and precedents through structured prompts. Cognitive load theory supports separating recall from organisation, and voice-first capture helps preserve the visual-manual channel during work (Sweller, 1988; Wickens, 2008). Outputs are presented as fragments rather than conclusions (DP1).

Structuring organises fragments into a coherent explanatory account through abduction (Peirce, 1997): inference to the best explanation using incomplete evidence. The sensemaking literature breaks this down into evidence foraging and schema fitting (Pirolli and Card, 2005). EASCI's counter-hypothesis requirement (DP7) operationalises the principle that competing narrative accounts must be evaluated against the same evidence

base (Bex et al., 2010). The quality gate requires: (i) the proposed explanation, (ii) at least one counterhypothesis with documented disposition, and (iii) evidence that would disprove the primary explanation.

Consolidation involves subjecting provisional KAs to peer validation, operationalising the social property of sensemaking (constructing shared interpretations of ambiguous situations) (Weick, 1995). The validator must provide a stated rationale for acceptance, modification, or rejection; dissenting views are recorded as part of provenance, not suppressed. AI is deliberately excluded from the validation role, for three independent reasons: the argument that expert judgment cannot be computationally reduced (Dreyfus and Dreyfus, 1986), Collins's finding that collective tacit knowledge cannot be held by any single agent (Collins, 2010), and the EU AI Act's human-oversight requirement.

Innovation governs validated KAs through periodic review, triggered review (when external events invalidate validation conditions), and explicit retirement (archiving KAs that no longer hold while preserving provenance records for audit). AI monitors usage patterns and flags KAs approaching review thresholds; human professionals decide whether to confirm, update, restrict, or retire.

4.2 Progressive Abstraction

The five stages progressively compress knowledge from rich, situated experience into structured, portable organisational memory. Experience and Articulation operate at the highest frequency, near the moment of action. Structuring and Consolidation operate at a lower frequency, integrating material across multiple capture episodes. Innovation operates at the lowest frequency, governing periodic revision of organisational knowledge bases. This multi-timescale structure parallels March's (1991) distinction between exploitation (refining existing knowledge through high-frequency operations) and exploration (discovering new possibilities through low-frequency reassessment).

4.3 Feedback Architecture

Designated feedback paths make SECI's cyclicity structurally explicit rather than metaphorical. They distinguish single-loop learning, which corrects actions within existing norms, from double-loop learning, which revises those norms (Argyris and Schön, 1978), and add deep-loop frame revision plus embodied micro-feedback. Figure 1 shows all eleven paths and makes cyclical knowledge creation auditable at the process level. Double-loop and deep-loop paths (F1, F7, F8, F9) modify organisational parameters, interpretive frames, hypothesis libraries, or validation norms; single-loop paths (F2-F5, F10-F11) adjust execution within existing norms. F5 is the only forward-directed single-loop path (Consolidation → Innovation), while the others return to earlier stages. F6 operates within Experience as the embodied precursor of the others. Low = months; Mid = days to weeks; Micro = within-episode.

The feedback paths also clarify how EASCI operationalises SECI beyond the forward decomposition of Externalisation (see also Figure 2). *Corrective Combination* occurs when explicit quality signals are combined with existing artefacts to produce revised artefacts, as in F2, F4, and F11. *Generative Combination* occurs when Innovation detects cross-KA patterns and updates hypothesis templates through F8. *Internalisation* occurs at several timescales: F1 returns validated procedural changes to practice conditions, F7 revises the interpretive frames practitioners bring to future articulation, and F9 changes the peer-validation norms through which later artefacts are assessed. The *Consolidation* gate is therefore not only a review checkpoint; it is the event where a provisional KO becomes socially explicit, while recorded peer feedback can be combined with the KO to produce a revised artefact.

5. Design Principles

Seven design principles outline rules and constraints embedded in EASCI's architecture, organised into three categories: *cognitive plausibility* (DP1-2), *social calibration* (DP3-4), and *temporal governance* with AI boundaries (DP5-7). Each addresses a failure mode that becomes more apparent under high-stress conditions rather than being purely theoretical. Figure 3 associates each principle with its corresponding failure mode and mechanism.

Cognitive plausibility DP1–DP2	Social calibration DP3–DP4	Temporal gov. & AI DP5–DP7
DP1 – Cognitive overload (E, A) <i>Risk:</i> Externalisation stalls under high cognitive load. <i>Control:</i> Accept fragmentary capture; separate recall from evaluation.	DP3 – Ungoverned reuse (S) <i>Risk:</i> Knowledge is applied outside its intended domain. <i>Control:</i> Require typed applicability conditions on each artefact.	DP5 – Knowledge decay (I) <i>Risk:</i> Outdated artefacts persist and drive current practice. <i>Control:</i> Periodic and triggered review with explicit retirement.
DP2 – Decontextualised use (E) <i>Risk:</i> Artefacts lose the constraints of their original context. <i>Control:</i> Capture operative constraints at point of action.	DP4 – Trust deficit (C) <i>Risk:</i> Peer validation collapses into rubber-stamping or silence. <i>Control:</i> Peer sign-off with recorded dissent and rationale.	DP6 – AI displaces judgement (All) <i>Risk:</i> GenAI suggestions override human responsibility. <i>Control:</i> Stage-bounded AI roles and an Evidential Authority Hierarchy.
		DP7 – Confirmation bias (S) <i>Risk:</i> Single narratives go unchallenged. <i>Control:</i> Mandatory counter-hypothesis with evidential grounding.

Figure 3: Design principles as structured knowledge-management risk controls

Cognitive plausibility (DP1-DP2). DP1 acknowledges that tacit knowledge is resistant to full verbalisation because expertise is cognitively compiled (Anderson, 1982). Fragments, provisional explanations, and incomplete analogies are regarded as valid contributions to knowledge. DP2 states that context must be captured as active constraints during action, as removing context disrupts the adaptive relationship between judgment and environment (Gigerenzer et al., 1999; Ericsson and Simon, 1993).

Social calibration (DP3-DP4). DP3 requires that each KA clearly state its reuse conditions to prevent inappropriate use outside the intended context. DP4 mandates explicit, socially grounded validation: peer review by practitioners with relevant domain experience, with dissenting opinions recorded. Trust functions in two ways: a primary barrier when absent, and a key enabler when present (Irbe and Ogunyemi, 2025).

Temporal governance and AI boundaries (DP5-DP7). DP5 supports revision, retirement, and intentional forgetting through three mechanisms (Kluge and Gronau, 2018). DP6 constrains GenAI to explicit stage boundaries: AI must not author KAs without human input, must not validate artefacts, must not autonomously retire knowledge, and must not suppress recorded dissent. An Evidential Authority Hierarchy ranks evidence sources from most trusted (operational traces, domain documents, expert testimony) to least trusted, with large language model (LLM) synthesis at the lowest level. DP7 requires every provisional explanation to be tested against at least one counterhypothesis, combining a Devil’s Advocate role with Socratic questioning (Nemeth et al., 2001; Walton, 1998).

6. GenAI Role Boundaries

Table 1 outlines GenAI permissions and restrictions at each stage. The counter-hypothesis requirement (DP7) aims to improve transparency by encouraging engagement with alternatives. PROV-O-aligned attribution structures specify who or what generated, used, transformed, or endorsed a KA; append-only or unalterable logging is an implementation-layer requirement, not a property proved by the formal model alone. The Consolidation gate is designed to prevent validation bypass by requiring human peer sign-off before a KA can advance; effective prevention still depends on interface-level enforcement. The Innovation stage manages drift and declining relevance through lifecycle metadata and review triggers. In high-stakes domains where decisions have legal or safety implications, these boundaries are regulatory requirements rather than mere design choices.

Table 1: GenAI role boundaries in EASCI stages

Stage	AI may	AI must not	Human responsibility
E	Capture metadata; transcribe	Interpret or prioritise experience	Identify what to externalise
A	Generate prompts; transcribe	Author content; complete explanations	Verbalise reasoning; author the KA
S	Suggest patterns; propose counterhypotheses	Determine validity; rank explanations	Select and endorse explanatory account
C	Surface related artefacts; flag conflicts	Validate; override dissent	Peer review; sign-off
I	Track usage; flag drift or declining relevance	Retire artefacts autonomously	Revision and retirement decisions

EASCI offers Process-Level Explainability (PLE) through five inspectable dimensions:

(PLE-1) *Provenance*: who articulated what, traced via PROV-O;

(PLE-2) *Rationale*: which alternatives were considered;

(PLE-3) *Review*: who reviewed and what dissent was recorded;

(PLE-4) *Reuse conditions*: boundaries under which knowledge applies; and

(PLE-5) *Evolution*: how and why knowledge has changed over revision cycles. Unlike model-level explainable AI, PLE functions independently of the underlying AI technology and directly addresses EU AI Act Art. 13 requirements for transparency.

6.1 Compliance Mapping

EASCI's architecture is designed to support compliance with the requirements of the EU AI Act. Human oversight (Art. 14) is structurally maintained: each stage transition involves a human decision point, and Consolidation necessitates documented peer review. The Evidential Authority Hierarchy meets Art. 14(4)(a) by restricting LLM-only evidence to Level 4. Technical documentation (Art. 12) is provided via the PROV-O attribution chain. Risk management (Art. 9(2)) is managed through the counter-hypothesis requirement and review triggers during the Innovation stage. Contestability (Art. 13) is supported by dissent preservation and the counter-hypothesis record; EASCI's intermediate artefacts support the architectural framework needed for challenging reasoning chains, as outlined by Almada (2019). When KAs include personal data, the Consolidation sign-off aligns with the meaningful human involvement criterion specified in GDPR Art. 22(3).

7. Evaluation Framework

Six falsifiable propositions test three claims: decomposition improves cognitive performance (P1), contextual anchoring improves artefact quality (P2, P4), and governance improves accountability (P3, P5, P6). P2 also bears on decomposition because it tests whether the separate Experience stage preserves contextual fidelity, but it is grouped here under contextual anchoring to avoid double-counting. Each proposition is operationalised for high-stress environments, where incomplete externalisation directly affects organisational decision-making.

P1 (Decomposition). EASCI's three-stage externalisation creates knowledge artefacts with completion rates and Artefact Traceability Scores (ATS) at least 1.2 times those of monolithic capture.

P2 (Context fidelity). Guided prompts based on Experience-stage metadata generate at least 1.5 times more contextual discriminations per artefact than open-ended requests.

P3 (Reviewability). Consolidated artefacts reach an inter-rater agreement of $\kappa \geq 0.70$ on the suitability of reuse, while unconsolidated artefacts reach $\kappa < 0.50$.

P4 (GenAI robustness). GenAI outputs from EASCI-produced artefacts show at least 30% lower rates of context mismatch and unsupported claims.

P5 (Cognitive load). Voice-first articulation with EASCI prompts results in lower NASA-TLX scores and higher completeness than text-based capture.

P6 (Lifecycle). EASCI's Innovation-stage mechanisms produce at least three times more versioning or retirement actions per review cycle, with at least 80% supported by documented rationale.

P1-P2 test the framework's core claim: if staged decomposition fails to produce better artefacts, the decomposition thesis requires a fundamental redesign.

P3-P6 evaluate peripheral mechanisms whose failure prompts targeted revisions of specific stages rather than the abandonment of the entire framework (Lakatos, 1970).

8. Discussion

EASCI does not introduce a new ontology of knowledge: it maintains SECI's tacit-explicit distinction and the understanding of knowledge conversion. The contribution operates within what Simon (1996) describes as "the sciences of the artificial": illustrating how knowledge is externalised, validated, and governed in practice. This situates EASCI within a design science epistemology (Hevner et al., 2004) that treats utility as the primary criterion of validity, rather than correspondence to a pre-existing theoretical model.

SECI treats externalisation as a single phase, assuming experts can articulate tacit knowledge through dialogue. EASCI argues that the three operations involved differ cognitively: *Articulation* is a recall task under high intrinsic load, *Structuring* is a hypothesis-generation task requiring metacognitive control, and *Consolidation* is a social-

judgment task whose validity depends on independent reviewers. Collapsing them into one phase produces either cognitive overload or social rubber-stamping. EASCI therefore assigns each operation an inspectable intermediate artefact with explicit provenance and reuse conditions. Where Nonaka's Ba assumes a shared context that persists over time, EASCI treats context as an artefact that must be actively constructed, recorded, and reused asynchronously.

Three conditions absent in 1995 now justify this extension: *GenAI-mediated co-production* requires explicit provenance and human-AI attribution; SECI offers no mechanism for this. *Regulatory accountability* (EU AI Act, LED) demands audit trails. *Empirical testability* requires falsifiable propositions with operationalised thresholds.

This addresses a gap that three decades of SECI research have not closed (Gourlay, 2006). EASCI's evaluation architecture specifies what to measure, the rejection threshold for each proposition, and the design of each test.

Two recent AI-extended SECI frameworks invite comparison. One *maps AI capabilities* onto each SECI mode descriptively, without specifying artefact requirements, or testable propositions (Zhang et al., 2025). The *GRAI framework* (Böhm and Durst, 2025) introduces a machine-agent layer without resolving the operational question: what must a practitioner produce at each stage, and who must sign off? EASCI differs in specifying typed artefacts, explicit gate criteria, and a role hierarchy in which AI suggestions cannot advance beyond Structuring without human grounding. Three non-claims deserve stating: the framework does not assert faithful representation of tacit knowledge; it does not claim every domain requires all five stages; and it treats GenAI as a bounded collaborator, not a bottleneck to minimise.

9. Boundary Conditions and Conclusion

Three requirements limit EASCI adoption. *BC1 (Domain complexity)*: EASCI applies where expertise depends on experience, is hard to articulate, and is difficult to formalise; the test is whether the knowledge would differ in another practitioner's hands given the same inputs. *BC2 (Organisational readiness)*: Consolidation assumes psychological safety and authorised challenge. *BC3 (Expert availability)*: Experience and Articulation depend on access to practitioners with tacit knowledge. A further limitation is empirical scope: EASCI was developed in EU law enforcement and fraud investigation, and has not yet been tested in healthcare, emergency response, or education.

Four domain-specific study designs represent the immediate next steps. First three specify each a Collins (2010) tacit knowledge type being tested (RTK, STK, CTK) and fourth is on lifecycle governance:

RTK (within-subjects controlled): in emergency medicine, P1 (decomposition) and P2 (context fidelity) could be tested by comparing post-incident debriefings with and without EASCI-structured prompts that measure the completeness and traceability of the resulting artefacts.

CTK (inter-rater reviewability): in a legal compliance team, P3 (reviewability) could be tested by having independent reviewers assess EASCI-consolidated KAs against informally documented procedures, using inter-rater agreement as the outcome measure.

STK (simulation-based transfer): in a professional development context in higher education, P5 (cognitive load) and P2 (context fidelity) could be assessed by comparing voice-first EASCI articulation with standard reflective journal entries, using NASA-TLX scores.

Lifecycle Governance: for any domain with a deployed knowledgebase, P6 (Lifecycle) can be tested by measuring the versioning and retirement action rate and markedly better documentation of lifecycle decisions.

These designs share the same propositions and rejection thresholds as the investigative case, ensuring cross-domain comparability without redesigning the evaluation architecture.

Three fidelity levels make adoption accessible regardless of technological maturity: low fidelity uses structured templates; medium fidelity integrates stage transitions into organisational workflows; high fidelity involves GenAI as a bounded collaborator with PROV-O attribution. Tacit knowledge externalisation remains the least operationalised element of knowledge management theory, thirty years after Nonaka and Takeuchi's foundational account (Nonaka and Takeuchi, 1995).

EASCI's core claim is that externalisation is not a single cognitive operation but three. Articulation reveals partial reasoning through structured prompts, Structuring organises fragments through abductive reasoning with mandatory counterhypotheses, and Consolidation subjects the result to peer validation. Experience captures situated context, and Innovation governs the lifecycle. Each operation produces a distinct artefact, passes

through a quality gate, and carries provenance metadata. Seven design principles bound GenAI to scaffolding roles rather than autonomous authority, and six falsifiable propositions provide an empirical agenda that SECI research has largely lacked. Testing those propositions across the four study designs outlined above, and in domains such as healthcare, education, and emergency response, is the next step. The current contribution is therefore a design whose broader generalisability remains a testable claim.

Acknowledgement

This study was co-funded by the European Union and the Estonian Ministry of Education and Research via the project TEM-TA26: Digital Transformation Through Lifelong Learning.

Ethics Declaration: The qualitative interviews discussed in this paper were carried out as part of a design science research project that received ethical approval from the relevant institutional review board. All participants gave informed consent. Interview data were anonymised and stored on encrypted, access-restricted systems. No personal data is included in this paper.

AI Declaration: AI tools (specifically, large language model assistants) were used during the preparation of this manuscript for language editing and formatting assistance. All substantive intellectual content, including the EASCI framework design, theoretical argumentation, empirical analysis, and evaluation architecture, was produced by the author. The author takes full responsibility for the content of this paper.

References

- Anderson, J.R. (1982) "Acquisition of cognitive skill", *Psychological Review*, Vol. 89, No. 4, pp. 369-406.
- Argyris, C. and Schön, D.A. (1978) *Organizational Learning: A Theory of Action Perspective*, Addison-Wesley, Reading, MA.
- Arksey, H. and O'Malley, L. (2005) "Scoping studies: towards a methodological framework", *International Journal of Social Research Methodology*, Vol. 8, No. 1, pp. 19-32.
- Barsalou, L.W. (2020) "Challenges and opportunities for grounding cognition", *Journal of Cognition*, Vol. 3, No. 1, p. 31.
- Bex, F.J. et al. (2010) "A hybrid formal theory of arguments, stories and criminal evidence", *Artificial Intelligence and Law*, Vol. 18, No. 2, pp. 123-152.
- Böhm, K. and Durst, S. (2026) "Knowledge management in the age of generative artificial intelligence - from SECI to GRAI", *VINE Journal of Information and Knowledge Management Systems*. Vol. 56, No. 1, pp. 106-201.
- Collins, H. (2010) *Tacit and Explicit Knowledge*, University of Chicago Press, Chicago.
- Costa, A.C.M.P. and Fantini, J.B.O. (2025) "Impact of generative artificial intelligence on knowledge management in software engineering: a systematic mapping study", *Proceedings of the 24th Brazilian Symposium on Software Quality (SBQS)*, São José dos Campos, pp. 55-66.
- Dewey, J. (1938) *Logic: The Theory of Inquiry*, Henry Holt and Company, New York.
- Dreyfus, H.L. and Dreyfus, S.E. (1986) *Mind over Machine: The Power of Human Intuition and Expertise in the Era of the Computer*, Free Press, New York.
- Dwivedi, Y.K., Kshetri, N., Hughes, L., et al. (2023) 'So what if ChatGPT wrote it?' Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI", *International Journal of Information Management*, Vol. 71, p. 102642.
- Eraut, M. (2004) "Informal learning in the workplace", *Studies in Continuing Education*, Vol. 26, No. 2, pp. 247-273.
- Ericsson, K.A. and Simon, H.A. (1993) *Protocol Analysis: Verbal Reports as Data*, revised edition, MIT Press, Cambridge, MA.
- European Parliament and Council of the European Union (2016) Directive (EU) 2016/680 (Law Enforcement Directive), *Official Journal of the European Union, L series*.
- European Parliament and Council of the European Union (2024) Regulation (EU) 2024/1689 (Artificial Intelligence Act), *Official Journal of the European Union, L series*.
- Gigerenzer, G., Todd, P.M. and the ABC Research Group (1999) *Simple Heuristics That Make Us Smart*, Oxford University Press.
- Gourlay, S. (2006) "Conceptualizing knowledge creation: a critique of Nonaka's theory", *Journal of Management Studies*, Vol. 43, No. 7, pp. 1415-1436.
- He, X., Antonczak, L. and Burger-Helmchen, T. (2026) "Revisiting Ba: how generative AI transforms knowledge creation", *Journal of Knowledge Management*. Vol. 30, No. 2, pp. 846-896.
- Hevner, A.R. et al. (2004) "Design science in information systems research", *MIS Quarterly*, Vol. 28, No. 1, pp. 75-105.
- Hinds, P.J., Patterson, M. and Pfeffer, J. (2001) "Bothered by abstraction: the effect of expertise on knowledge transfer and subsequent novice performance", *Journal of Applied Psychology*, Vol. 86, No. 6, p. 1232.
- Irbe, I. (2025) "EASCI design science: mathematical formalization, knowledge phase portrait, and machine-checked proofs", *in preparation*.
- Irbe, I. and Ogunyemi, A.A. (2025) "Designing for the unspoken: a work-in-progress on tacit knowledge transfer in high-stress public institutions", *Proceedings of the 36th Annual Conference of the European Association of Cognitive Ergonomics (ECCE 2025)*, ACM.
- Klein, G. (1998) *Sources of Power: How People Make Decisions*, MIT Press.

- Kluge, A. and Gronau, N. (2018) "Intentional forgetting in organizations: the importance of eliminating retrieval cues for implementing new routines", *Frontiers in Psychology*, Vol. 9, p. 51.
- Lakatos, I. (1970) "Falsification and the methodology of scientific research programmes", in Lakatos, I. and Musgrave, A. (eds.), *Criticism and the Growth of Knowledge*, Cambridge University Press, pp. 91-196.
- March, J.G. (1991) 'Exploration and exploitation in organizational learning', *Organization Science*, Vol. 2, No. 1, pp. 71-87.
- Mayer, R.C., Davis, J.H. and Schoorman, F.D. (1995) "An integrative model of organizational trust", *Academy of Management Review*, Vol. 20, No. 3, pp. 709-734.
- Nemeth, C.J. et al. (2001) "Devil's advocate versus authentic dissent: stimulating quantity and quality", *European Journal of Social Psychology*, Vol. 31, No. 6, pp. 707-720.
- Nonaka, I. and Konno, N. (1998) "The concept of 'Ba': building a foundation for knowledge creation", *California Management Review*, Vol. 40, No. 3, pp. 40-54.
- Nonaka, I. and Takeuchi, H. (1995) *The Knowledge-Creating Company: How Japanese Companies Create the Dynamics of Innovation*, Oxford University Press, New York.
- Peirce, C.S. (1997) *Pragmatism as a Principle and Method of Right Thinking: The 1903 Harvard Lectures on Pragmatism*, edited by P.A. Turrissi, State University of New York Press, Albany.
- Pirolli, P. and Card, S. (2005) "The sensemaking process and leverage points for analyst technology as identified through cognitive task analysis", *Proceedings of the International Conference on Intelligence Analysis*, Vol. 5, No. 1, pp. 2-4.
- Polanyi, M. (1966) *The Tacit Dimension*, University of Chicago Press, Chicago.
- Simon, H.A. (1996) *The Sciences of the Artificial*, 3rd edition, MIT Press, Cambridge, MA.
- Sweller, J. (1988) "Cognitive load during problem solving: effects on learning", *Cognitive Science*, Vol. 12, No. 2, pp. 257-285.
- Szulanski, G. (1996) "Exploring internal stickiness: impediments to the transfer of best practice within the firm", *Strategic Management Journal*, Vol. 17, No. S2, pp. 27-43.
- Tsoukas, H. (2005). Do we really understand tacit knowledge. *Managing knowledge: an essential reader*, pp. 107-127.
- Walton, D.N. (1998) *The New Dialectic: Conversational Contexts of Argument*, University of Toronto Press, Toronto.
- Weick, K.E. (1995) *Sensemaking in Organizations*, Sage Publications, Thousand Oaks, CA.
- Wickens, C.D. (2008) "Multiple resources and mental workload", *Human Factors*, Vol. 50, No. 3, pp. 449-455.
- Zhang, W., Zhang, W., Daim, T., and Yalcin, H. (2025) "AI challenges conventional knowledge management: light the way for reframing SECI and Ba theory", *Journal of Knowledge Management*. Vol. 29, No. 5, pp. 1618-1654.