

## **SPECTRA HUMAN-CENTERED LEARNING FRAMEWORK: A SOCIO-TECHNICAL ARCHITECTURE FOR AI-SUPPORTED HIGHER EDUCATION**

MARCO SPECTRA DE APRENDIZAJE CENTRADO EN EL SER HUMANO: UNA ARQUITECTURA SOCIOTÉCNICA PARA LA EDUCACIÓN SUPERIOR APOYADA POR IA

**Mario Guadalupe López Ayala** mario.lopez@uadeo.mx  
Universidad Autónoma de Occidente, México.  
ORCID: <https://orcid.org/0000-0001-5377-6257>

**Cómo citar este artículo / Citation:** López Ayala, M. G. (2026). Spectra Human-Centered Learning Framework: A Socio-Technical Architecture for AI-Supported Higher Education. *Revista Científica Arbitrada de la Fundación MenteClara*, Vol. 11 (441).  
DOI: <https://doi.org/10.32351/rca.v11.441>

**Copyright:** © 2026 RCAFMC. Este artículo de acceso abierto es distribuido bajo los términos de la licencia [Creative Commons Attribution 4.0 International License \(CC BY 4.0\)](https://creativecommons.org/licenses/by/4.0/).  
Recibido: 22/09/2026. Aceptado: 25/09/2026. Publicación online: 29/09/2026.

**Conflicto de intereses:** Ninguno que declarar.

### **Abstract**

Generative artificial intelligence (GenAI) can produce persuasive university work without revealing who performed or checked the reasoning behind it. This conceptual article develops an AI-specific elaboration of the previously disseminated Spectra Human-Centered Learning Framework. Its three foundational conditions, six pedagogical phases and five original propositions are retained, not presented as new inventions. A purposive, critical synthesis distinguishes the functions of the pedagogical cycle and examines how AI assistance, intellectual responsibility, institutional safeguards and contextual constraints may intersect with them. The resulting architecture connects real-world input, conceptual framing, engagement, production, assessment and reflective transfer. The inherited propositions remain verbatim; task-level interpretations, candidate indicators and conditions that might challenge them are added for future research. Augmentation, dependency and

substitution describe provisional distributions of work between learners and systems, not validated learner types. A targeted comparison with adjacent frameworks identifies overlaps and the specific scope of this elaboration without claiming global novelty or superiority. No student data were collected, no intervention was implemented and no systematic literature search was conducted. Whether the configuration supports independently defensible learning, accessible assessment and transfer across contexts remains an empirical question.

## **Resumen**

La inteligencia artificial generativa (IAGen) permite elaborar productos universitarios convincentes, aunque su calidad no permite atribuir automáticamente al estudiante las decisiones que contienen. Este artículo propone una elaboración del Spectra Human-Centered Learning Framework orientada a entornos con IA y reconoce expresamente la formulación anterior: tres condiciones fundacionales, seis fases pedagógicas y cinco proposiciones que aquí se conservan. Mediante síntesis conceptual crítica de bibliografía seleccionada intencionalmente, se distinguen las funciones del ciclo formativo y se examina cómo pueden intervenir la asistencia de IA, la responsabilidad intelectual y las condiciones institucionales. El resultado es una arquitectura teórica que relaciona insumos del mundo real, encuadre conceptual, participación, producción, evaluación y transferencia; no es una intervención implementada. Las proposiciones heredadas se mantienen literalmente y se acompañan de interpretaciones y posibles criterios de refutación para investigaciones futuras. La ampliación, la dependencia y la sustitución designan relaciones provisionales entre personas y sistemas a escala de tarea, no perfiles validados del alumnado. El contraste selectivo con otras propuestas delimita el aporte sin reclamar eficacia o prioridad general. No se recogieron datos de estudiantes ni se realizó una revisión sistemática. Queda por examinar si la organización propuesta favorece razonamiento independiente, accesibilidad y transferencia en contextos distintos.

**Keywords:** AI governance; artificial intelligence; higher education; human-centered learning; socio-technical systems

**Palabras Claves:** aprendizaje centrado en el ser humano; educación superior; gobernanza de IA; inteligencia artificial; sistemas sociotécnicos

## 1. Introduction

Generative artificial intelligence (GenAI) has entered ordinary university work: students use it to locate evidence, draft explanations, calculate, translate and obtain feedback. Yet the finished artifact leaves a practical question unanswered. Which decisions belong to the student, which were delegated, and which could the student defend without assistance? The issue concerns what an assessment can legitimately demonstrate, not merely whether a tool was permitted (Bond et al., 2024; Kasneci et al., 2023).

Earlier research on AI in higher education catalogued applications while identifying underrepresentation of educators and pedagogical perspectives (Zawacki-Richter et al., 2019). A later review of 67 empirical studies reports context-dependent critical- and creative-thinking outcomes associated with structured GenAI use, not a general causal advantage of using AI (Li et al., 2026). A survey of 319 knowledge workers concerns reported cognitive effort and checking practices; because it is neither a university-student experiment nor a longitudinal cognition study, it cannot establish that GenAI use causes cognitive decline among students (Lee et al., 2025).

Established educational research offers complementary mechanisms. Constructive alignment connects intended outcomes, activities and assessment (Biggs, 1996); active-learning evidence foregrounds meaningful student participation (Freeman et al., 2014); authentic assessment relates evaluated performance to disciplinary practice (Gulikers et al., 2004). Multimedia-learning theory specifies constraints on representation and cognitive processing (Mayer, 2020), self-regulated-learning research addresses monitoring and strategy adjustment (Panadero, 2017), and transfer research cautions that successful task completion does not guarantee performance in another context (Barnett & Ceci, 2002).

These educational traditions are not displaced by GenAI. Instead, the same system can now enter several moments of instruction, from the first encounter with evidence to the revision of a final answer. Human-centered AI draws attention to consequential human control (Shneiderman, 2020); socio-technical research locates that control within organizations, practices and infrastructures (Baxter & Sommerville, 2011). Their intersection is important, but it does not by itself tell an instructor what to teach, what to assess or where intellectual responsibility should reside.

Spectra HCL offers a pedagogical starting point. Its earlier formulation already specified three foundational conditions, six phases and five conceptual propositions

(López Ayala, 2026). This paper asks what happens to that architecture when GenAI is available throughout the process. It does not claim to originate the inherited components. The present elaboration instead distinguishes transversal assistance from teaching phases, locates responsibility at the level of individual tasks, and introduces human and institutional safeguards alongside contextual constraints. Whether these additions constitute a useful analytical distinction must remain open to scrutiny; adding layers to a diagram is not itself a contribution to learning.

The earlier full-text version and its figure were publicly disseminated through Mendeley Data as a conceptual research output (López Ayala, 2026); this antecedent is cited rather than described as a peer-reviewed journal article. The present manuscript retains the earlier pedagogical sequence and original propositions, revises the research problem toward AI-supported learning, and elaborates responsibilities and safeguards across the cycle. The relation between the two works is disclosed to distinguish continuity from subsequent conceptual development.

The proposed contribution is a specification that can be challenged: it identifies pedagogical functions, assigns possible responsibilities, and makes conditions of implementation visible. This has implications for accessible and accountable course design, although a coherent description cannot substitute for evidence of learning, equity or agency.

The objective is to articulate an inspectable, AI-specific version of the inherited Spectra architecture and identify claims that could be examined rather than assumed. Three questions organize the inquiry: (RQ1) How can established learning mechanisms be configured as an AI-supported university learning architecture? (RQ2) Which observable distributions of intellectual responsibility and safeguards might distinguish assistance from inappropriate delegation in a specified task? (RQ3) What disciplinary, linguistic, infrastructural and governance conditions should be documented when the architecture is implemented and tested? The third question concerns contextual implementation of Spectra, not a procedure for transferring an intervention between institutions.

## **2. Methodology**

### ***2.1. Research design and conceptual approach***

The study uses a critical integrative conceptual synthesis and model-development design. Jaakkola (2020) distinguishes theory synthesis and model development as approaches to conceptual articles: the former integrates knowledge from partially separated traditions and the latter makes constructs and their proposed relationships explicit. These orientations guide the integration of educational and socio-technical mechanisms into the proposed learning architecture. No comprehensive or standardized Jaakkola protocol is claimed.

Guidance on integrative reviews (Torraco, 2005; Snyder, 2019) informs the cross-tradition reading; Jabareen (2009) provides a vocabulary for differentiating constructs and their relationships. These sources orient the argument rather than certify a completed protocol. No study of students, courses or universities was undertaken, and no empirical validation of the architecture is claimed.

### ***2.2. Fixed documentary corpus, selection and evidence roles***

The predecessor manuscript developed nine theoretical strands and cited 17 references, including Jaakkola (2020). The present inquiry adds purposively selected work on AI-supported education, conceptual methods, assessment and governance, together with targeted attention to neighboring frameworks. Inclusion was governed by argumentative relevance: a source had to help define a mechanism, expose a design tension, characterize human control or make an implementation condition visible. This is a documentary corpus selected for conceptual work, not a sample intended to estimate the prevalence of findings.

Foundational theoretical studies inform the proposed definitions and functional relationships; empirical reports support only appropriately bounded statements about the populations and designs actually studied; policy and institutional sources motivate implementation questions rather than establish learning outcomes. The original literature assembly was purposive, not exhaustive or statistically representative. No dated database strategy, search strings, retrieval totals, exclusion log, independent coding, PRISMA screening, or formal risk-of-bias appraisal exists for that historical work. Later source additions and retrospective mapping do not transform it into a systematic review or a prospectively documented evidence

synthesis. The complete reference list identifies all literature cited in this version, whereas Table 1 provides selected, not exhaustive, source-to-function links.

### ***2.3. Functional differentiation and mapping rules***

The analysis proceeded through four interpretative operations: isolate a source's relevant mechanism, distinguish it from nearby constructs, place it provisionally within a phase or cross-cutting condition, and examine what AI assistance changes in the allocation of work. Constructive alignment concerns the relation among intended outcomes, activities and assessment (Biggs, 1996). Authentic assessment, by contrast, asks what kind of performance evidence a task elicits (Gulikers et al., 2004). The distinction matters even though the constructs meet in classroom practice: an aligned activity can still yield insufficient evidence of independent competence.

Phase assignments follow a functional rule: source encounter and representation belong primarily to Phase 1; disciplinary framing and interpretative categories to Phase 2; active examination and questioning to Phase 3; learner-authored demonstrable work to Phase 4; appraisal of competent performance to Phase 5; and application to meaningfully changed conditions to Phase 6. The distinction is analytical, not a claim that the phases are non-overlapping psychological factors or that every lesson must follow six chronological steps. Reflection can operate during production, and assessment can prompt a return to framing. An alternative five-phase or seven-phase configuration is a legitimate competitor, not a violation of the theory.

AI assistance, source verification, metacognitive monitoring and transparency may recur across phases. They are treated as transversal functions or safeguards, not as a new step that learners must reach after reflective transfer. Infrastructure, language, culture, access and data governance similarly condition implementation without becoming pedagogical phases. This allocation is a modeling decision, not the only defensible one.

### ***2.4. Construction of the architecture and theoretical traceability***

Three inherited foundations organize the synthesis: human-centered orientation, socio-technical awareness and cross-disciplinary applicability. Contextual

constraints shape how the latter can be realized. The six learning phases distinguish primary instructional purposes, but not six sealed psychological compartments. Assessment may send a learner back to framing; an unfamiliar application may initiate another cycle. GenAI can assist with exploration, comparison, organization or feedback at several points and is therefore specified transversally in the present text, not as a seventh phase. The accompanying safeguards concern provenance, monitoring, transparency, judgment, accountability and contestability. Their inclusion signals obligations for design; it does not establish that they have reduced risk in practice.

The configuration was subjected to an argumentative check: each phase should have a distinct primary purpose, a traceable conceptual basis, a plausible relationship to neighboring phases, and some prospect of empirical discrimination. Table 1 makes selected correspondences and possible indicators explicit. It is not a full extraction register and records no measurements. A retrospective editorial worksheet can help another reader inspect present choices, but it cannot recreate decisions never documented during the original work.

## ***2.5. Derivation of conceptual propositions***

Five conceptual propositions are retained verbatim from the predecessor manuscript (López Ayala, 2026). Their functions are to specify possible relationships among framing, applied engagement, anticipatory assessment, reflective transfer and socio-technical orientation. The present article does not manufacture a new set of five findings; it examines how each inherited proposition might be interpreted and potentially disconfirmed in AI-supported settings. Proposed indicators and comparison conditions are prospective operationalizations developed here, not measures used in the original deposit or data collected for this article. Their wording and numbering remain P1–P5 in Section 4.2.

Augmentation, dependency and substitution are proposed as task-level descriptors of where intellectual work occurs. They should not be inferred from the number of prompts a student enters. Nor do they form an inevitable trajectory or a psychometric scale. The meaningful question is narrower: relative to a stated learning objective, which operations can the learner still perform, explain and defend?

## 2.6. Reproducibility boundary and methodological limitations

The reference list, the declared functional mapping rules, selected correspondences in Table 1 and a separately prepared retrospective worksheet support inspection of the current interpretation but do not reconstruct the full contemporaneous decisions behind the predecessor. The historical search and coding process cannot be replicated from an absent original search log. A third party can challenge particular source-to-component assignments, propose alternative arrangements and investigate whether the five inherited propositions follow from the present AI-specific elaboration; this is conceptual inspectability, not empirical replicability or proof of a unique six-phase solution.

Contextual heterogeneity in cited studies may generate hypotheses about implementation; it does not establish educational effectiveness across countries or causal geopolitical effects. Subsequent investigations should prespecify the AI tools and versions, degree of access, instructional objectives, comparator conditions, reasoning and source-verification outcomes, accessible independent assessments, transfer tasks and contextual measures. An appropriately designed negative case would also examine learning tasks where AI-free work or a reduced phase arrangement performs as well as, or better than, the proposed configuration.

Table 1. Theoretical traceability of the proposed architecture

| Source or tradition    | Construct and Spectra position            | Candidate future indicator                |
|------------------------|-------------------------------------------|-------------------------------------------|
| Mayer (2020)           | Multimodal cognitive processing → Phase 1 | Comprehension and processing demands      |
| Biggs (1996)           | Constructive alignment → Phase 2          | Conceptual justification of a task        |
| Freeman et al. (2014)  | Active engagement → Phase 3               | Reasoning quality during engagement       |
| Panadero (2017)        | Self-regulation → Phase 4 and safeguards  | Planning, monitoring and independent work |
| Gulikers et al. (2004) | Authenticity → Phase 5                    | Oral defense and performance evidence     |
| Barnett & Ceci (2002)  | Transfer → Phase 6                        | Performance on a novel academic problem   |

| Source or tradition                | Construct and Spectra position           | Candidate future indicator               |
|------------------------------------|------------------------------------------|------------------------------------------|
| Shneiderman (2020)                 | Human control → Transversal safeguards   | Source checking and decision authority   |
| Baxter & Sommerville (2011)        | Socio-technical design → Foundations     | Institutional implementation differences |
| OECD (2026); Hamadeh & Amin (2025) | Governance and international constraints | Access, privacy and local feasibility    |

Note. Author’s conceptual mapping based on sources named in the first column. Indicators are proposed for future validation and were not measured in this study.

3. Results

The result is a conceptual elaboration of the antecedent Spectra HCL framework. The three foundational conditions, six phases and five propositions belong to the earlier formulation (López Ayala, 2026). The present article locates GenAI assistance across those phases, describes how intellectual responsibility could be distributed within tasks, and identifies safeguards and contextual conditions for examination. Table 1 sets out selected theoretical connections; Table 2 distinguishes continuity from the AI-specific contribution. Figure 1 reproduces the inherited pedagogical core. It does not depict the additional transversal AI functions or institutional safeguards, which are specified in Sections 3.2–3.4. Neither the figure nor the tables report measured educational effects.

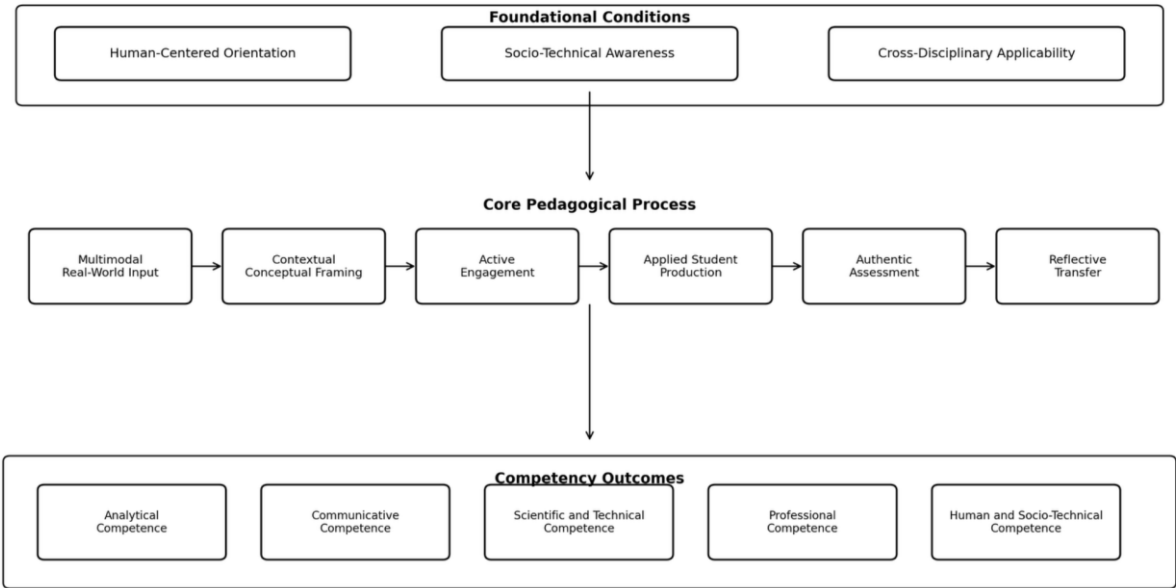
Table 2. Genealogy and scope of the revised Spectra contribution

| Component        | Earlier Spectra formulation                                | Present AI-specific elaboration                                   |
|------------------|------------------------------------------------------------|-------------------------------------------------------------------|
| Pedagogical core | Six functions from input to reflective transfer            | Functions retained; treated as iterative with candidate AI roles  |
| Foundations      | Human-centered, socio-technical and contextual orientation | Responsibilities and governance specified within those conditions |

| Component          | Earlier Spectra formulation                                                                                                | Present AI-specific elaboration                                                                               |
|--------------------|----------------------------------------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------------------------|
| AI position        | Digital mediation considered broadly                                                                                       | GenAI assistance across six retained phases; no new instructional phase is proposed.                          |
| Human–AI relations | Learner agency and meaningful learning emphasized                                                                          | Task-level augmentation, dependency and substitution proposed as unvalidated heuristics                       |
| Propositions       | Five propositions: framing, engagement and production, anticipatory assessment, transfer, and socio-technical orientation. | Retained verbatim; AI interpretations, candidate indicators and counterevidence are discussed in Section 4.2. |

Note. “Earlier” denotes the public May 2026 Mendeley Data Version 1 (DOI: 10.17632/fvfcw7898h.1). Retained elements should not be misrepresented as introduced for the first time in this AI-specific elaboration.

Figure 1. Relational Architecture of the Spectra Human-Centered Learning Framework



Note. This figure reproduces the foundational conditions, six-phase pedagogical sequence and intended competency outcomes of the earlier Spectra formulation (López Ayala, 2026). The arrows express an instructional design relationship, not an experimentally established causal pathway. Although drawn sequentially, the phases may be revisited. The transversal AI assistance, safeguards and implementation conditions elaborated in the present article are discussed in Sections 3.2–3.4 and are not represented in this inherited figure. The competencies are intended outcomes, not findings measured in this study.

### **3.1. Foundational Conditions and Six Learning Phases**

The diagram retains three foundational conditions. Human-centered orientation positions learners as agents who can interpret and challenge outputs; socio-technical awareness situates instruction within platforms, policies and organizational responsibilities. Cross-disciplinary applicability makes a more modest claim: the pedagogical functions may be examined in different fields, while language, resources and disciplinary forms of evidence will still require attention. These are design premises, not validated predictors of student achievement (Baxter & Sommerville, 2011; Shneiderman, 2020).

Phase 1, multimodal real-world input, connects learning to an authentic phenomenon, dataset, primary source, case or observation. Material should be selected for relevance and designed to avoid extraneous processing; AI-generated summaries must be distinguishable from original evidence (Mayer, 2020). Phase 2, contextual conceptual framing, establishes terms, disciplinary explanations, theoretical alternatives and assessment criteria. AI may compare explanations, but its fluency does not establish contextual validity or source reliability (Biggs, 1996; Kasneci et al., 2023).

Phase 3, active engagement, requires analysis, discussion, comparison, questioning or independent problem framing rather than mere exposure. AI can serve as a critic or simulation partner, provided learners examine its assumptions (Freeman et al., 2014; Kahu, 2013). Phase 4, applied student production, converts intellectual activity into a defensible academic artifact. The learner must explain why evidence was selected and which decisions were made, whether AI supported drafting, calculation or revision. Autonomous motivation and self-regulatory monitoring remain relevant to this design (Ryan & Deci, 2020; Panadero, 2017).

Phase 5, authentic assessment, requests evidence of competence in meaningful disciplinary tasks, which may include source provenance, work traces and oral, written, practical or accessible alternative justification. Artifact quality alone cannot prove independent understanding (Gulikers et al., 2004; Ates, 2026). Phase 6,

reflective transfer, asks learners to adapt their knowledge to a genuinely changed question, context or data condition rather than restate an earlier response (Barnett & Ceci, 2002). Each phase identifies a primary design purpose, not a mandatory classroom ritual. A well-designed microtask may combine engagement and production, and feedback may start a new cycle; the full six-phase configuration is proposed at the level of an instructional sequence or module. Its incremental contribution relative to a reduced alternative must be tested.

The phases are distinguished by what an instructor needs to know: the evidence encountered, the meaning assigned to it, the inquiry performed, the work produced, the basis of judgment and the prospect of transfer. A discussion may involve production; reflection need not wait until the end. This overlap is not a defect to be removed by definition. It creates an empirical problem: a shorter sequence might achieve comparable independent learning, and Spectra has not yet shown that all six phases are necessary.

### ***3.2. AI Support and Safeguards***

GenAI is placed across the learning cycle because assistance with exploration, translation, comparison, simulation and feedback can occur at several stages. Assigning it a seventh phase would obscure this fact. Operational proficiency with a tool is not equivalent to the reflective AI literacy required to recognize uncertainty or question provenance (Lintner, 2024). The model suggests a distinction; it does not assume that a student who can explain a prompt can also evaluate the evidentiary status of its answer.

Human-facing safeguards include provenance checking, metacognitive monitoring, interpretation of context and accountable decisions. Institutional safeguards include transparent permissions, training, data governance, accessible tool alternatives, identified responsibility for consequential decisions and practical means to contest them. They are distributed obligations rather than a demand that students independently resolve all system risks (OECD, 2026). Naming safeguards is a design proposal: no risk-reduction effect has been demonstrated here.

### **3.3. Augmentation, Dependency and Substitution**

The proposed human–AI relationship continuum is best examined at the level of a particular task. In augmentation, support coexists with the learner’s retained ability to explain and check the target decision. Dependency would require evidence of reduced independent flexibility after assistance is withdrawn, assessed against a suitable baseline or comparator; tool frequency alone is insufficient. Substitution occurs when the system performs an operation designated as a learning objective and the learner cannot reconstruct or defend it under an appropriate independent assessment. For a non-target operation, delegation may be acceptable. One student may display different arrangements across tasks.

These are provisional analytical descriptions, not validated personality categories, psychometric dimensions or stages through which users inevitably pass. Evidence offers reasons to investigate them but not to treat them as measured Spectra constructs (Li et al., 2026). In a 68-student quasi-experiment, Xu et al. (2025) found advantages in self-regulated learning with metacognitive support but no significant between-group achievement difference. Lee et al. (2025) surveyed knowledge workers about perceived cognitive effort and checking. That evidence cannot establish causal cognitive decline among university students. The preference for retained intellectual agency is a normative design choice; its educational consequences require a separate test.

### **3.4. Institutional, Cultural and Geopolitical Conditions**

The learning architecture is not assumed to function within a technologically uniform world. Differences in infrastructure, access, institutional policy, language resources and data governance may alter which AI supports can be provided and which safeguards are feasible. The geopolitical dimension refers specifically to international distributions of technological capability, platform dependence and regulatory divergence; it should not be conflated with every cultural or institutional difference (Hamadeh & Amin, 2025; OECD, 2026).

Regional sources provide reasons to test contextual variation, not a basis for deterministic regional claims. Valentini (2026) documents institutional adoption and governance heterogeneity in Latin America and the Caribbean; a separate exploratory study of university students in Chihuahua concerns its own sampled participants and cannot represent all Mexican universities (Chávez Márquez & De los Ríos Chávez,

2025). Neither source isolates geopolitics as a cause of individual cognition. Student characteristics, teaching practices, discipline, platform access and institution-level policy must be measured at the appropriate analytic level.

A shared conceptual language need not imply interchangeable institutional settings. A practice may be feasible in one course and inappropriate in another because of privacy rules, equipment, language coverage or assessment design. Whether any of these conditions changes the relationship between assistance and learning must be investigated rather than inferred from national labels.

## **4. Discussion**

### ***4.1. Academic Application, Positioning and Future Validation***

A hypothetical assignment illustrates the distinction. Students examine primary sources, challenge competing explanations and produce an analysis that must later be applied to a changed case. GenAI could supply counterarguments, translation or feedback; the task would still need independently defined opportunities for source checking and justification, with accessible alternatives to any oral defense. This is an illustration of design choices, not an implemented course or evidence that Spectra improved performance.

The surrounding literature limits any easy claim to novelty. Oliveira et al. (2025) developed DRIVE to examine student–GenAI interactions and reported an association with essay scores in a particular setting; that association cannot validate Spectra. Chapman and Dalgarno (2026) connect assessment authenticity, evidence provenance and governance. Uden and Hwang (2026) address effortful processing, reflection and scaffolding through the LEARN framework; Vendrell and Johnston (2026) likewise develop critical-thinking design principles, while Ates (2026) experimentally compares feedback arrangements. These works already address agency, process evidence and responsible AI support. Spectra makes a narrower proposal: to locate related concerns across its previously formulated six-phase learning cycle while distinguishing task responsibilities from institutional conditions. The comparison is selective. It cannot establish priority or superiority over all alternatives.

## **4.2. Research propositions and potential disconfirmation**

The five propositions below retain the original wording and numbering of the predecessor (López Ayala, 2026). Their accompanying interpretations belong to this AI-specific article. They describe relationships worth investigating, not five results obtained from students. Each test would need an independently specified outcome and a credible comparison; contrary findings would matter as much as apparent support.

**Proposition 1. The educational value of multimodal input depends on the quality of contextual conceptual framing that accompanies it.**

AI-specific interpretation. When learners encounter AI-generated summaries or multimodal material, conceptual framing can expose provenance, disciplinary categories and uncertainty rather than merely supply the next prompt. Framing quality and subsequent comprehension should be assessed separately. If comparable tasks with and without such framing yield no difference—or the expected relation is reversed—P1 would require revision. A fluent generated summary is not evidence that its reader understood it.

**Proposition 2. Active engagement produces deeper learning when it is connected to applied output rather than limited to participation alone.**

The distinction in P2 becomes delicate when a system conducts most of the discussion. Engagement would need to involve decisions the learner can explain; applied production offers one way to make that responsibility inspectable. A comparison of engagement-only with engagement-plus-production conditions should measure independent reasoning, not just the polish of AI-assisted artifacts. Equivalent or poorer learning in the latter would challenge P2.

**Proposition 3. Authentic assessment is strongest when it is anticipated at the design stage of inputs, framing, and tasks, not added only at the end of instruction.**

For P3, assessment cannot be treated as a final inspection detached from what the task allowed. Criteria should specify permitted assistance, expected evidence and accountability before work begins. Process traces and independent explanations may be informative, but alternatives to oral defense must remain accessible and artifact

quality should be evaluated separately. If advance alignment adds nothing to a distinct measure of learning, this proposed relationship would be weakened.

**Proposition 4. Reflective transfer increases when learning activities explicitly connect disciplinary content to professional, organizational, technological, or social contexts.**

Transfer poses a stricter question than whether students can prompt the same system again. An assessment would require an independently examined task with meaningful contextual change. Holding earlier instruction and access constant, researchers could vary explicit links to professional or social settings and then assess independent performance elsewhere. Null effects, adverse effects or gains limited to nearly identical tasks would narrow the claim made in P4 (Barnett & Ceci, 2002).

**Proposition 5. In technology-rich higher education, pedagogical coherence requires a socio-technical orientation that preserves learner agency and meaning rather than treating technology as pedagogically neutral.**

P5 places responsibility across more than the student–tool interaction. Learner decision rights, accessible alternatives, institutional authority, data arrangements and linguistic conditions all matter, but their presence cannot be used circularly to define the very coherence being tested. The relevant features and outcomes would need independent measures at task, course and institutional levels. Comparable or contrary patterns could challenge stronger interpretations of P5. A country label, by itself, explains neither a governance mechanism nor a cognitive outcome.

An empirical program could begin with independent experts assigning the disclosed sources to functional categories and recording disagreements. Task-level definitions and accessible measures would then need piloting. Only after these steps would comparisons between Spectra-informed modules and credible alternative designs—including AI-free arrangements where appropriate—be interpretable. Potential outcomes include independent reasoning, provenance checks, artifact quality, explanation and performance on novel tasks; P5 additionally requires distinct learner-, course- and institution-level measures. Reports should specify system versions, instructor preparation, task demands, prior competence, language, infrastructure and permitted assistance. None of these investigations was conducted for the present article.

### **4.3. Limitations**

The central limitation is evidentiary. This is a conceptual synthesis, not a systematic review, intervention trial, scale validation or population estimate. The predecessor was assembled purposively, without a dated search record or complete contemporaneous coding ledger; the later worksheet cannot turn absent historical records into prospective documentation. The six phases and five propositions preceded the present paper. Functional boundaries, safeguarding effects and the task-level human–AI descriptors have not been empirically validated.

Bibliographic identity was checked against available records, but the full text of every source was not examined for every manuscript claim. Model capability, language coverage, discipline, accessibility and institutional governance also change across settings and over time. The targeted comparison with neighboring frameworks cannot establish global originality or superiority. Alternative phase arrangements and AI-free instruction remain valid comparator conditions. Observational reports are not evidence of caused cognitive deterioration; a regional sample cannot stand for global higher education.

## **5. Conclusions**

Spectra HCL retains its previously disseminated three foundations, six phases and five propositions. This article places GenAI assistance across that inherited sequence and identifies tasks in which responsibility, verification and institutional safeguards may become salient. Its principal outcome is a researchable conceptual configuration—not evidence that the configuration works better than other designs. The original figure makes the pedagogical core visible; the AI-specific elements are stated and differentiated in the text.

The original propositions relate framing to input, engagement to production, assessment to design, transfer to contextual connections, and coherence to socio-technical responsibility. The AI-specific interpretations propose ways to challenge those relationships, not five new empirical findings. Augmentation, dependency and substitution remain descriptions of work allocation within tasks. Their value will depend on whether independent measures can distinguish them from prior competence, task difficulty and permitted tool use.

What remains uncertain is substantial. Future comparisons must examine independent reasoning, defensible work traces, authentic and accessible assessment,

and transfer under genuinely changed conditions. Five- and seven-phase alternatives deserve examination alongside the inherited six-phase sequence. The human-centered premise is to preserve occasions for learners to understand, verify and contest the knowledge they present; whether this particular architecture helps them do so is not yet known.

**Declaration on generative AI assistance.** ChatGPT (OpenAI) was used to assist with conceptual organization, methodological critique, drafting, language editing, and manuscript formatting. The author retains responsibility for the scholarly claims, source verification, and submitted content.

## References

- Ates, H. (2026). Human-centered GenAI feedback design in higher education: A multisite experiment on direct, reflective, and hybrid approaches to scientific argumentation. *International Journal of Educational Technology in Higher Education*, 23, Article 38. <https://doi.org/10.1186/s41239-026-00614-9>
- Barnett, S. M., & Ceci, S. J. (2002). When and where do we apply what we learn? A taxonomy for far transfer. *Psychological Bulletin*, 128(4), 612–637. <https://doi.org/10.1037/0033-2909.128.4.612>
- Baxter, G., & Sommerville, I. (2011). Socio-technical systems: From design methods to systems engineering. *Interacting with Computers*, 23(1), 4–17. <https://doi.org/10.1016/j.intcom.2010.07.003>
- Biggs, J. (1996). Enhancing teaching through constructive alignment. *Higher Education*, 32, 347–364. <https://doi.org/10.1007/BF00138871>
- Bond, M., Khosravi, H., De Laat, M., Bergdahl, N., Negrea, V., Oxley, E., Pham, P., Chong, S. W., & Siemens, G. (2024). A meta systematic review of artificial intelligence in higher education: A call for increased ethics, collaboration, and rigour. *International Journal of Educational Technology in Higher Education*, 21, Article 4. <https://doi.org/10.1186/s41239-023-00436-z>
- Chapman, B. L., & Dalgarno, B. (2026). Beyond the ban: A theoretical framework for integrating generative AI in assessment. *Policy Futures in Education*, 24(5), 670–682. <https://doi.org/10.1177/14782103251411725>
- Chávez Márquez, I. L., & De los Ríos Chávez, H. J. (2025). Uso personal y académico de inteligencia artificial en estudiantes universitarios: estudio exploratorio. *Apertura*, 17(1), 54–69. <https://doi.org/10.32870/ap.v17n1.2604>
- Freeman, S., Eddy, S. L., McDonough, M., Smith, M. K., Okoroafor, N., Jordt, H., & Wenderoth, M. P. (2014). Active learning increases student performance in science, engineering, and mathematics. *Proceedings of the National Academy of Sciences*, 111(23), 8410–8415. <https://doi.org/10.1073/pnas.1319030111>
- Gulikers, J. T. M., Bastiaens, T. J., & Kirschner, P. A. (2004). A five-dimensional framework for authentic assessment. *Educational Technology Research and Development*, 52(3), 67–86. <https://doi.org/10.1007/BF02504676>
- Hamadeh, S., & Amin, H. (2025). AI, education and digital sovereignty. *Frontiers in Education*, 10, Article 1677727. <https://doi.org/10.3389/feduc.2025.1677727>
- Jaakkola, E. (2020). Designing conceptual articles: Four approaches. *AMS Review*, 10(1–2), 18–26. <https://doi.org/10.1007/s13162-020-00161-0>
- Jabareen, Y. (2009). Building a conceptual framework: Philosophy, definitions, and procedure. *International Journal of Qualitative Methods*, 8(4), 49–62. <https://doi.org/10.1177/160940690900800406>
- Kahu, E. R. (2013). Framing student engagement in higher education. *Studies in Higher Education*, 38(5), 758–773. <https://doi.org/10.1080/03075079.2011.598505>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, Article 102274. <https://doi.org/10.1016/j.lindif.2023.102274>

- Lee, H.-P. H., Sarkar, A., Tankelevitch, L., Drosos, I., Rintel, S., Banks, R., & Wilson, N. (2025). The impact of generative AI on critical thinking: Self-reported reductions in cognitive effort and confidence effects from a survey of knowledge workers. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, Article 1121, 1–22. <https://doi.org/10.1145/3706598.3713778>
- Li, C., Cui, H., & Hagedorn, L. S. (2026). The cognitive impact of ChatGPT in higher education: A systematic review of critical and creative thinking outcomes. *Computers and Education: Artificial Intelligence*, 10, Article 100571. <https://doi.org/10.1016/j.caeai.2026.100571>
- Lintner, T. (2024). A systematic review of AI literacy scales. *npj Science of Learning*, 9, Article 50. <https://doi.org/10.1038/s41539-024-00264-4>
- López Ayala, M. (2026). Spectra Human-Centered Learning Framework: A human-centered and socio-technical model for cross-disciplinary university learning (Version 1) [Conceptual manuscript and accompanying figure]. Mendeley Data. <https://doi.org/10.17632/fvfcw7898h.1>
- Mayer, R. E. (2020). *Multimedia learning* (3rd ed.). Cambridge University Press. <https://doi.org/10.1017/9781316941355>
- OECD. (2026). Policies supporting responsible and systematic GenAI adoption in higher education (OECD Education Spotlights, No. 23). OECD Publishing. <https://doi.org/10.1787/c4e5621f-en>
- Oliveira, M., Zednik, C., Bombaerts, G., Sadowski, B., & Conijn, R. (2025). Assessing students' DRIVE: A framework to evaluate learning through interactions with generative AI. *Computers and Education: Artificial Intelligence*, 9, Article 100497. <https://doi.org/10.1016/j.caeai.2025.100497>
- Panadero, E. (2017). A review of self-regulated learning: Six models and four directions for research. *Frontiers in Psychology*, 8, Article 422. <https://doi.org/10.3389/fpsyg.2017.00422>
- Ryan, R. M., & Deci, E. L. (2020). Intrinsic and extrinsic motivation from a self-determination theory perspective: Definitions, theory, practices, and future directions. *Contemporary Educational Psychology*, 61, Article 101860. <https://doi.org/10.1016/j.cedpsych.2020.101860>
- Shneiderman, B. (2020). Human-centered artificial intelligence: Reliable, safe & trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495–504. <https://doi.org/10.1080/10447318.2020.1741118>
- Snyder, H. (2019). Literature review as a research methodology: An overview and guidelines. *Journal of Business Research*, 104, 333–339. <https://doi.org/10.1016/j.jbusres.2019.07.039>
- Torraco, R. J. (2005). Writing integrative literature reviews: Guidelines and examples. *Human Resource Development Review*, 4(3), 356–367. <https://doi.org/10.1177/1534484305278283>
- Uden, L., & Hwang, G.-J. (2026). The LEARN framework for responsible use of generative AI in education: A neuroscience-informed model for problem-based learning. *Journal of Computers in Education*. Advance online publication. <https://doi.org/10.1007/s40692-026-00398-x>
- Valentini, A. (2026). Adopción y gobernanza de la IA en la educación superior de América Latina y el Caribe: resultados de una encuesta regional. *Revista Educación Superior y Sociedad*, 37(2), 372–397. <https://doi.org/10.54674/ess.v37i2.1278>
- Vendrell, M., & Johnston, S.-K. (2026). Scaffolding critical thinking with generative AI: Design principles for integrating large language models in higher education. *Computers and Education: Artificial Intelligence*, 10, Article 100572. <https://doi.org/10.1016/j.caeai.2026.100572>
- Xu, X., Qiao, L., Cheng, N., Liu, H., & Zhao, W. (2025). Enhancing self-regulated learning and learning experience in generative AI environments: The critical role of metacognitive support. *British Journal of Educational Technology*, 56(5), 1842–1863. <https://doi.org/10.1111/bjet.13599>

Zawacki-Richter, O., Marín, V. I., Bond, M., & Gouverneur, F. (2019). Systematic review of research on artificial intelligence applications in higher education—Where are the educators? *International Journal of Educational Technology in Higher Education*, 16, Article 39.  
<https://doi.org/10.1186/s41239-019-0171-0>