

Review

Agentic AI as Co-Comparatist: Autonomous Literary Intelligence and The Posthuman Turn in World Literature Studies

Viraj P. Tathavadekar^{1,*}, Aleksandar Erceg², Pallavi Priya Silveira¹

¹Research Scholar, Symbiosis International (Deemed University), Pune, Maharashtra, India

²Faculty of Economics and Business in Osijek, Josip Juraj Strossmayer University of Osijek, Osijek, Croatia

*Corresponding author: Viraj P. Tathavadekar, virajtatu@gmail.com

Abstract

This conceptual paper asks a narrow question with wide consequences. When an artificial intelligence system stops behaving like a retrieval instrument and starts decomposing goals, selecting its own tools, and issuing interpretive statements without a prompt for each move, does anything in comparative literature change, and if so what, and under what conditions? The discipline still tends to treat AI as a search or translation aid, while the technical literature has moved to systems that act across several steps with little supervision. The paper does not argue that machines read, understand, or enjoy literature. It defends a deliberately limited proposition: that such a system can take on the functional role of a co-comparatist, proposing intertextual relations across corpora that human scholars must then evaluate, and that this functional shift is enough to unsettle the anthropocentric foundation of comparative method. The method is theory adaptation applied to a structured corpus of ninety recent sources on agentic AI, set beside foundational scholarship in world literature, translation studies, and digital humanities that carries no date restriction, and read through posthumanist theory and agential realism. The paper separates three things: machine pattern detection, machine-assisted interpretation, and autonomous scholarly interpretation. It offers a claim-level evidential audit, a four-layer governance cycle in which marginalised communities participate from corpus design onward, a worked example tested under favourable and adverse conditions, and eight conceptual propositions with proposed measures. The reverse risk is weighed equally: a corpus skewed toward Western canons would automate Eurocentric hierarchy at scale.

Keywords

Agentic AI, Comparative literature, Posthumanism, Decolonial digital humanities, Distributed agency, Conceptual framework

Article History

Received: 20 July 2026

Revised: 02 September 2026

Accepted: 08 September 2026

Available Online: 23 September 2026

Copyright

© 2026 by the authors. This article is published by the Cultech Publishing Sdn. Bhd. under the terms of the Creative Commons Attribution 4.0 International License (CC BY 4.0): <https://creativecommons.org/licenses/by/4.0/>

1. Introduction

The field of comparative literature takes its title from its crossing of borders. It passes between languages, countries, eras, and media, and makes such crossing its technique, rather than merely an embellishment. One border it has yet to cross is the border between human and machine reading, and this border is now under attack from an unexpected angle. This is something worth saying right from the start, since the subsequent argument rests on getting the scale of the assault correctly sized up, neither exaggerating it into a revolution nor reducing it to a mere tool.

Most current scholarship treats computation as an extension of the reading self. The case of attempts to bring together short stories from various traditions into one curriculum reveals the comparative urge at its finest, and the analysis involved remains completely done by humans [1]. Machine translation has been shown to have gained significant progress in terms of accessing languages in comparison but acknowledges that idioms, registers, and culturally-loaded references fare very poorly [2]. In cross-cultural language learning systems, the machine serves as scaffolding that recedes once the competence is attained [3]. It follows, then, that all these pieces carry the same presupposition: the machine prepares the materials for analysis by the scholar.

That assumption is doing more harm than good. Posthumanist literacy scholarship has argued for some years that everyday AI participates in meaning rather than merely conveying it, and that reading and writing are better theorised as distributed practices [4,5]. Work on conversational agents in language immersion sharpens the point, finding that the abstract, context-free register of digital agents sits awkwardly beside the embodied ethos that posthumanist applied linguistics describes [6]. Meanwhile the technical literature has moved on. Systematic reviews describe agentic systems that decompose goals, select tools, and act across several steps with little supervision [7,8], and information systems research now treats agentic architectures as an organisational form rather than a software feature [9]. Perspectives on agentic architecture set out the planning and tool-use components in detail [10]. The distance between how comparative literature talks about AI and what these systems now do is widening.

A small body of work has begun to close the distance without quite reaching comparative literature. One study examines credence and attribution in the context of generating literary meaning through a generative system, and asks to whom interpretative credit is owed if the machine performs the reading [11]. The agential realism perspective holds that since the system of production now includes non-human elements, responsibility for the meaning generated cannot reside solely with the human agent [12]. Research on autonomy in philosophical terms poses the related problem of whether machine-generated output can be considered knowledge in the first place [13], while skepticisms provide the cautionary note that the framework currently in place might not be adequate to interpretation's needs [14]. Heritage studies alert us to the way that generative systems are already changing our understanding of cultural authenticity [15]. None of these addresses comparative literature directly, however, nor asks whether the autonomous system can make a comparative judgment, which is the definition of the field.

This paper addresses that omission with a claim that it is careful not to overstate. It is the proposition that an agentic system can serve as a co-comparatist, being able to formulate intertextual hypotheses, test them on multi-language databases, and produce claims for interpretation that humans must consider. The term co-comparatist is used in this paper as a functional and metaphorical designation of such a role rather than a statement that the machine can read, comprehend, and have scholarly accountability. There are two reasons why this question is timely rather than theoretical. First, studies on algorithmic cultural transmission have shown that AI systems already influence whether images and narratives are distributed throughout the world [16]. Second, reviews of AI in language education have indicated how fast such technologies get to be used as infrastructure once they prove to be useful [17], while a bibliometric survey of AI literature demonstrates how far behind criticism falls behind adoption [18].

The two research questions guiding this paper are conceptual.

RQ1: under what conditions will the output of an agentic system be elevated from pattern detection to a defensible comparative claim, and where will the scholarly responsibility lie in such a case?

RQ2: what form of governance structure would allow machine engagement to expand world literature without recolonising it?

There are three main contributions. The first contribution is the use of posthumanist literary theory and agentic systems to make the connection explicit that has not been made so far in the literature. The second contribution is the proposal of a governance cycle where the disciplinary specificity, along with its divergence from a human-in-the-loop cycle, is explicitly stated. The third contribution is the provision of evidential weight of every argument in this paper that separates demonstrable claims from analogical, speculative, and normative arguments. Section 2 introduces the literature. Section 3 introduces the conceptual approach and the evidentiary audit. Section 4 builds the construct, the cycle, an example of implementation, and eight propositions. Section 5 is about implications and limitations. Section 6 concludes.

2. Literature Review

2.1 Distributed Agency in Posthumanist Theory

Traditional literary scholarship rests on a premise it rarely states: that the human reader is the only entity capable of literary understanding. Posthumanist critique has questioned that premise for more than twenty years, treating agency as something distributed across networks of persons, texts, technologies, and institutions rather than lodged in a single mind [4,5]. The lineage matters, because the idea is older than the machines that now test it. N. Katherine Hayles [19] reframed the human as an information-processing entity continuous with its media, Rosi Braidotti [20] recast subjectivity as relational and more-than-human, and Bruno Latour's [21] insistence that non-human actors carry agency in a network gave the position its working vocabulary.

Critical work on urban AI extends the same argument to intelligence and autonomy, treating them as relational properties of assemblages rather than attributes of individuals [22], and accounts contrasting artificial with collective intelligence reach a compatible conclusion, that distributed cognitive systems have produced insight no single member held [23]. Reading these together supplies the theoretical warrant the paper needs: if agency was never purely human to begin with, a machine that proposes a relation is not an intruder into interpretation but a new participant in an assemblage that was always plural.

The strength of this approach for the argument made here is also its weakness. Distributed agency is an easy concept to invoke and difficult to define. Saying that agency is spread across a network does not by itself tell us how much of it sits with the corpus designer, how much with the model, how much with the validating scholar, and how much with the communities whose literatures are at stake. The agential realist literature insists that this distribution be made concrete, tying meaning and responsibility to the specific configuration of the apparatus that produces a reading [12]. A conceptual paper that leans on distributed agency therefore owes the reader an account of the distribution, not merely an appeal to the idea. Section 4 supplies that account.

2.2 Comparative Literature, Digital Humanities, and the Translation Question

Comparative reading is a cognitive achievement built over years. A scholar acquires several languages, internalises the conventions those languages carry, and learns to hold two texts closely enough that a structural echo becomes visible. The capacity is real and narrow at once, since few researchers work confidently across more than a handful of traditions, and the comparative map of world literature has therefore always been drawn from a small number of vantage points. The discipline has argued about exactly this limit for two decades. David Damrosch [24] defined world literature as writing that circulates and gains in translation, which makes the reader's linguistic reach constitutive of what counts as world literature at all. Franco Moretti [25] pressed the point to its provocation, arguing that no scholar can read the tens of thousands of novels the world has produced and proposing distant reading as the necessary substitute for the close reading a single career allows.

Pascale Casanova [26] mapped an unequal world literary space in which a few metropolitan centres confer value on the periphery, and Emily Apter [27] countered that the very ideal of frictionless world literature suppresses the untranslatable, the resistance that translation cannot dissolve. These are not settled questions, and the machine enters precisely here, into a field already divided over how far the reach of comparison can be extended without distorting what it compares. Curricular work that brings distant traditions into one frame shows both the reach and the labour of this practice [1].

The field of computational literary studies and distant reading has held out the promise of expanding the window, while translation studies has cautioned that any such expansion done in a mechanical way inevitably warps what is observed through it. Studies evaluating machine translation weigh both its gains and its losses [2]. Studies focusing on the role of AI in the teaching of translation see the machine as an intercultural mediator rather than as a simple channel, drawing on cultural schema theory and intercultural pragmatics [28]. Systems for cross-cultural language learning are concerned less with translation than with modelling culturally specific strategies [3]. Studies of AI-mediated cultural imagery show algorithmic systems already governing visibility at scale [16], and dataset work on cultural heritage question answering shows the same systems being pointed at non-Western material such as iconographies of Durga [29].

The politics of that pointing is the crux. Cross-domain generative work on ethnic design demonstrates how readily a model flattens culturally specific form into whatever its training distribution favours [30], and decolonial scholarship on the integration of African traditional knowledge into AI systems names the epistemic redress that would be required to resist such flattening [31]. Comparative analysis of how AI itself is framed in Chinese versus Western media shows that even the discourse about these systems is geographically skewed [32]. Reframings of digital literacy [33] and critiques

of algorithmic learning that ask whether automation erodes learner autonomy [34] round out a picture in which the tool is never culturally neutral.

The computational turn sharpened the same argument rather than resolving it. Distant reading and computational literary studies promised to widen the aperture, and the field's own critics supplied the cautions. Katherine Bode [35] challenged the field to account for the constructed, edited nature of the archives it mines rather than treating a corpus as given. Ted Underwood [36] showed what macro-scale modelling of literary history can reveal that close reading cannot, while conceding how much depends on the corpus's composition. Roopika Risam [37] and others working on postcolonial and decolonial digital humanities pressed the harder question of whose texts are digitised in the first place, arguing that a digital archive inherits and can amplify the exclusions of the print one. That line joins an older postcolonial critique.

Edward Said's [38] account of how the metropolitan canon narrated its others, Gayatri Spivak's [39] warning that the subaltern is spoken for rather than heard, and translation-studies work from Lawrence Venuti [40] on the translator's invisibility and Susan Bassnett [41] on translation as cultural rewriting all bear directly on a system that proposes cross-cultural comparisons from a skewed corpus. The technical sources this paper cites on translation, cultural imagery, and decolonial data [2,3,16,28-34] are read against this scholarship, not in place of it. The point of naming the debate is to locate the machine inside it: an agentic system that recommends European comparators for African or South Asian texts is not committing a novel error but automating a hierarchy the discipline has spent forty years trying to dismantle.

2.3 Agentic AI and What It Can Be Shown to Do

The word agentic needs a working definition before it can carry any weight in a literary argument. The definition of an agentic system in the technical literature is "one which breaks down a high-level task into sub-tasks, makes decisions about using tools, retains state through stages and acts with minimal human intervention" [7-10]. The infrastructure of large language models provides this kind of ability [42], and federated systems can extend it beyond distributed data [43]. None of these traits constitute autonomy in the philosophical sense of intentionality and consciousness, and this paper makes sure to separate technical autonomy from these richer concepts. What agentic systems demonstrably do is narrower than the vocabulary suggests and still consequential: they act across steps, and the direction of some of those steps is set by the system rather than by a person. Explainability research qualifies the picture further, since a system that acts without exposing why it acted poses obvious problems for scholarship [44], and work on prompt design shows how sensitive outputs remain to the exact framing of a request [45]. Agent-based generative models applied to land-use decisions illustrate the pattern of self-directed intermediate steps in a concrete setting [46].

2.4 The Bridge, and the Ethical Literature Around It

The gap this paper addresses is not that comparative literature ignores AI, but that it has no account of the specific transfer from autonomous multi-step behaviour to literary agency. That transfer is exactly what scrutiny should resist granting for free, and the paper treats it as something to be argued rather than assumed. The surrounding ethical literature marks the stakes. Work on the alignment problem shows how hard it is to make autonomous systems pursue the goals their principals intend [47]. Non-Western ethical frameworks, such as trusteeship ethics, extend the normative lexicon to include more than the Euro-American default position [48]. The literature on AI in risk-intensive areas like mental health care demonstrates the kind of governance required for autonomous systems when there are significant costs associated with failure [49]. This literature does not provide grounds for transferring to literary agency, but it establishes the conditions that such agency must satisfy.

2.5 Theoretical Basis

In the framework outlined above, AI is recast so that it is more than a tool while still remaining one. If, in fact, interpretation is an achievement in which many actors play a part [4,5,22], and if responsibility is determined by how the apparatus is configured and not by the type of the elements that make up that apparatus [12], then a system which establishes a justified comparison relationship has made a contribution to knowledge even though none of its parts understand that relationship. The contribution is incomplete by definition, and the missing piece, which involves judgment of an aesthetic and cultural sort, is the one supplied by human scholars.

3. Methodology

3.1 Research Design and Conceptual Research Questions

This is a conceptual paper that uses theory adaptation: a theory is carried into a setting where one of its assumptions no longer holds, and the resulting friction generates new constructs. The assumption under strain is anthropocentrism, the premise that literary interpretation is a human activity alone. The procedure follows the pattern common to review-based theoretical work in adjacent fields, where secondary sources replace primary data and the test of the output is conceptual adequacy rather than statistical fit [18,50,51]. A point of method matters here, because a reviewer rightly pressed it: the recency restriction described in Section 3.2 applies only to the technical literature on agentic AI, where the capability under discussion has no earlier counterpart. The disciplinary and theoretical literature carries no date limit. Foundational work in comparative and world literature, translation studies, and digital humanities enters the synthesis on the strength of its argument, not its publication year, and the older debates named in Section 2 are treated as load-bearing rather than as background. Figure 1 sets out the three-phase design. The two research questions from the introduction, RQ1 on the move from pattern to claim and RQ2 on governance, guide the synthesis throughout.

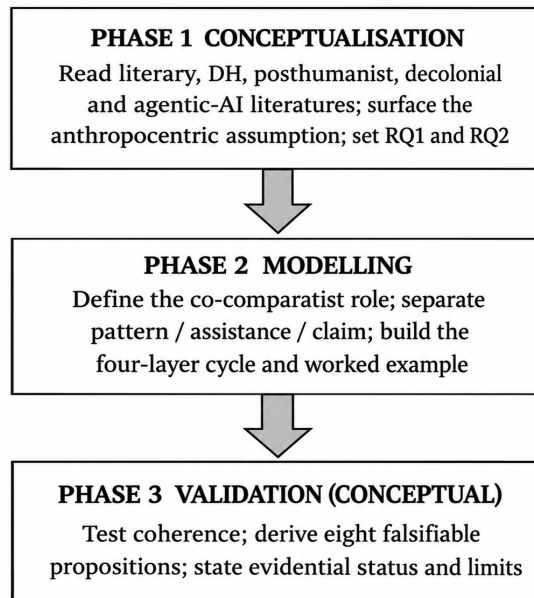


Figure 1. Conceptual research design.

Source: Authors' own elaboration.

3.2 Corpus Identification and Screening

The corpus was assembled through a structured search aimed at conceptual sufficiency rather than exhaustive coverage, and the search is reported here in enough detail to be repeated. Three databases were queried: Scopus, Web of Science, and the ACL Anthology, the last included to reach computational-linguistics work that the first two index unevenly. Searches were run between March and May 2026, with a final verification sweep in June 2026. The search string combined three concept blocks with Boolean operators: (agentic OR autonomous OR "large language model" OR "generative AI" OR "AI agent") AND ("comparative literature" OR "world literature" OR "distant reading" OR "digital humanities" OR translation OR "literary interpretation" OR canon) AND (posthuman* OR decolonial OR hermeneutic* OR "AI ethics"). For the technical block on agentic systems, records were restricted to 2021 onward, since systems of the relevant kind did not exist earlier. For the disciplinary and theoretical blocks, no date restriction was applied, and foundational sources were added by backward citation chaining from the retrieved articles. Screening proceeded in three passes: de-duplication by DOI; title and abstract screening against relevance to interpretive agency, comparison, or corpus governance; and full-text assessment. Records were excluded when AI appeared only as an application domain with no bearing on interpretation, comparison, authorship, or governance, and when a source duplicated a stronger one already retained.

Conceptual sufficiency was judged reached when additional records stopped introducing new themes or new disciplinary positions, a saturation criterion standard in review-based synthesis. Ninety sources were retained. Figure 2 reports the full identification, screening, and inclusion counts in PRISMA form. The counts are those of the present search and should be re-run by anyone replicating the corpus. Bibliometric and systematic-review studies informed the screening logic [52-56].

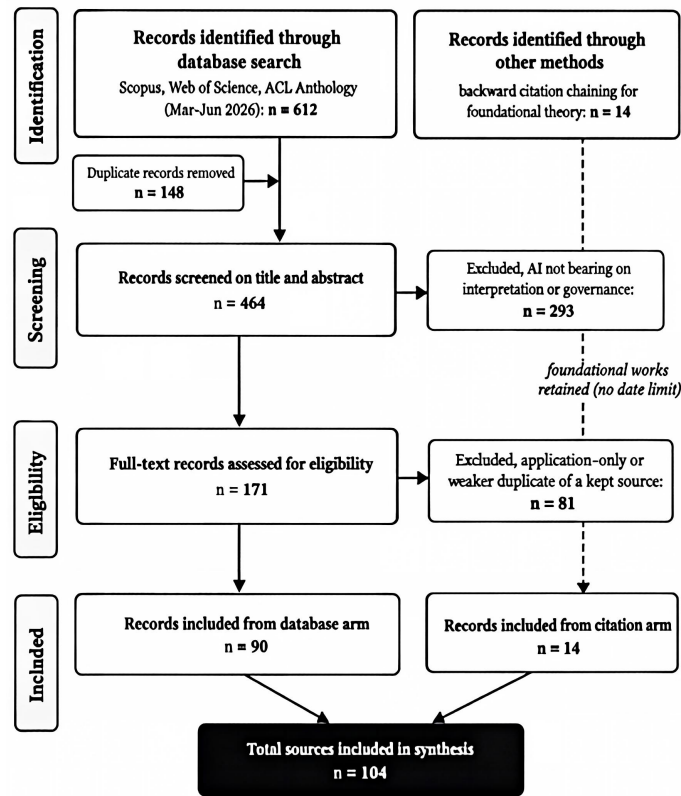


Figure 2. PRISMA-style flow of corpus identification, screening, and inclusion.

Source: Authors' own elaboration

3.3 Evidential Audit: Separating Direct Evidence from Analogy

A reviewer objected, fairly, that assigning references to broad evidential categories does not show whether any single claim is properly supported. This revision answers that objection in two ways. First, the most remote sources have been demoted. Studies of wine consumption, novel foods, hotel service, hospital layout generation, and autonomous vehicles appear now only where a specific mechanism is genuinely shared, and never as support for a literary conclusion; several that earlier served only to show that AI is spreading have been cut back to a single illustrative mention or removed from the argument's load path entirely. Second, and more importantly, Table 1 no longer sorts references into bins. It maps individual claims to their supporting sources and records, for each, whether the support is direct disciplinary evidence, a labelled cross-domain analogy, adoption evidence, or technical background. A reader can therefore check the warrant behind any load-bearing sentence rather than trusting a category label. The governing rule is unchanged and now enforced claim by claim: every analogy is marked as an analogy in the prose, and no literary conclusion rests on an analogy alone.

Table 1. Evidential audit: The role each source group plays in the argument.

Load-bearing claim (abbreviated)	Supporting source(s)	Evidential status
Agency is distributed across an interpretive assemblage	[4,5,12,22]	Direct disciplinary
World-literature reach is limited by a scholar's languages	[1,24,27]	Direct disciplinary
Machine translation gains access but distorts idiom and register	[2,3,28]	Direct disciplinary
Algorithmic systems already govern cultural visibility at scale	[16,29,30]	Direct disciplinary
Skewed corpora reproduce Eurocentric hierarchy	[30-34]	Direct disciplinary
An agentic system decomposes goals and acts across steps	[7-10]	Technical background
A system can maintain and update state across long sequences	[57-60]	Cross-domain analogy
Heterogeneous inputs can be fused into one analysis	[43,61-63]	Cross-domain analogy
AI moves from novelty to infrastructure once useful	[18,52,64,65]	Adoption evidence
Cross-cultural validation designs can check machine output	[66]	Technical background
Explainability and prompt sensitivity bound machine output	[43-45]	Technical background

Source: Authors' own elaboration.

3.4 Analytical Procedure

Analysis followed an inductive thematic sequence, and the coding is described here so that the synthesis can be assessed rather than taken on trust. From each source, statements were extracted where they bore on one of three concerns: the direction of interpretive agency, the durability of human authority, or the governance of corpora. Extracted statements were assigned open codes, and the codes were then grouped by shared analytical function rather than by shared topic, so that a claim about translation error and a claim about hallucinated attribution could sit under one code about the reliability of machine output.

Two rounds of coding were run, and code labels were revised only where a second reading would plausibly have assigned a statement differently. Three themes resulted: the status of machine output as pattern, assistance, or claim; the distribution of agency and responsibility; and the conditions of decolonial governance. Figure 3 traces the path from statement to theme. The classification in Table 1 was carried out on the same statements, so that evidential status and thematic role were fixed together rather than in separate passes.

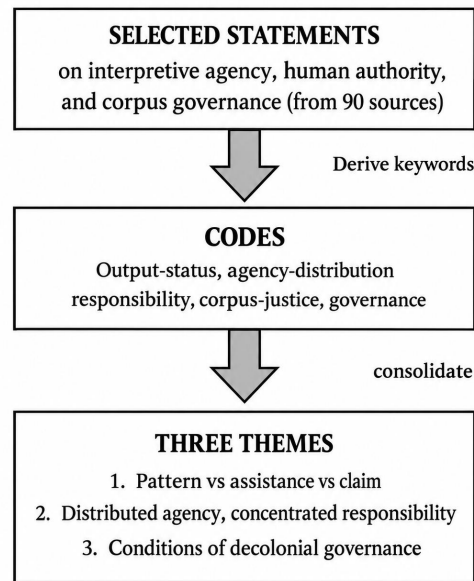


Figure 3. Inductive synthesis from statements to themes.

Source: Authors' own elaboration.

3.5 Conceptual Model Development and Validity

The model was built by treating each theme as a design layer and specifying determinants at the level at which they could in principle be observed, so corpus balance is expressed through documented composition rather than through stated intent. Propositions were written to be falsifiable against evidence that studies could later gather. Content validity rests on deriving the construct from two established bodies of theory rather than from intuition; construct validity is served by defining the co-comparatist role through observable behaviour rather than through claims about machine mind; and the limits of the exercise are stated in Section 5. Cross-cultural validation work offers a template for the empirical studies that would follow [66].

4. Theory Development

4.1 Three Things that Must be Kept Apart

A system that reports a recurring structure across a corpus is doing pattern detection. A system whose output a scholar uses to reach an interpretation is providing machine-assisted interpretation. A system whose proposal itself has the form of a comparative claim, one that can be right or wrong and that a scholar must accept, refine, or reject, is doing something closer to autonomous scholarly interpretation, though the responsibility for the claim never leaves the human side. Table 2 lays out the three, their evidential status, and the boundary conditions that separate them. The distinction matters because most of what agentic systems demonstrably do today sits in the first two columns. Pattern detection is well evidenced, in cultural imagery [16], ethnic design generalisation [30], attention-mapped model behaviour [44], and, by explicit analogy only, in signal and target recognition tasks in unrelated fields [67-69]. Machine-assisted interpretation is evidenced in translation and language pedagogy [2,3,17,28,70], in cross-cultural learning platforms [71], and in media substitution studies [72]. The third column, autonomous scholarly interpretation, is where the paper's claim lives, and it is advanced as a conditional possibility rather than a demonstrated fact.

Table 2. Three modes of machine involvement in comparative reading, and the boundary conditions that separate them.

Mode	What the machine does	Evidential status	Boundary condition
Pattern detection	Reports recurring structure across a corpus	Demonstrated [16,30,44]; analogy only [67-69]	Becomes assistance once a scholar builds an interpretation on it
Machine-assisted interpretation	Supplies material a scholar reason from	Demonstrated [2,3,17,28,70-72]	Becomes autonomous once the output itself has claim form
Autonomous scholarly interpretation	Proposes a comparative claim that can be right or wrong	Conditional possibility, not demonstrated [13,14]	Never carries responsibility: that stays with the validator

Source: Authors' own elaboration.

4.2 What Agentic Systems can be Shown to Do, Inferred to Do, and Merely Imagined to Do

Demonstrated: agentic systems break down goals, choose tools, and act through steps [7-10]; they process multilingual content in amounts greater than a research group [73]; they discover structure across huge corpora [16]. Derived by analogy to different domains, and noted as such: the ability of a system to maintain and modify its own state across many steps, a feature established in building control [57], micro-grid control [58], and autonomous navigation [59,60], mentioned here solely as an example that state-maintenance is possible, not as literature on the matter; and that heterogeneous inputs can be fused together, established in multimodal driver and sensor systems [61-63,74] and in federated model situations [43], once again as analogy.

Hypothetical: a single tradition could be described in a few hours, or that a scholar with a single language could work seriously across twenty. These are possibilities the technology gestures toward, not achievements, and the paper labels them so. Normative: everything in Section 4.4 about how such systems should be governed, which is a claim about what ought to happen, not about what does. This ordering directly answers the objection that hypothetical capabilities were presented as established ones. The scale and reach that earlier drafts asserted are here demoted to the speculative tier, and the load-bearing claims rest only on the demonstrated tier and on analogies that are marked as analogies.

4.3 The Co-Comparatist as Functional Role, and Where Agency Sits

The sharpest challenge to this label concerns the verb originates. If the system is still operating within the context of a corpus, a task, a prompt, a retrieval mechanism, and a validation framework that human beings construct, then describing the system as something more than a tool may be a mere re-description. The answer lies not in arguing but in making the difference visible. An action can be deemed an example of autonomy sufficient to merit the designation co-comparatist only if four criteria are satisfied simultaneously and verifiably through the system's own logs:

First, internal goal decomposition: the system establishes at least one sub-goal that was not specified in the prompt, such as which pairs of traditions to compare next.

Second, an unprompted comparison: The particular intertextual relationship retrieved did not have a name, an exemplar, or any prompting within the prompt.

Third, the form of output in the form of a claim: The output is a claim that can be evaluated as either true or false, accompanied by evidence.

Fourth, a substantially contingent trajectory: a re-run based on the same prompt but using a different retrieval seed produces a different solution, and hence the system, not the prompt, determined the trajectory.

Wherever any of the four conditions is not met, the system is engaging in interpretation aided by machines and is referred to as such. Whenever all four conditions are satisfied, the label denotes the functional role, not consciousness nor responsibility, but something more than mere retrieval. The conditions are specifically behavioral in order that one can deny the label on the basis of the log.

The co-comparatist claim can be stated precisely with the distinctions in place. The system is not an independent epistemic agent, because it does not own the claims it makes; it is not merely a tool, because it originates proposals rather than only executing instructions; it is a delegated analytical agent operating inside a distributed interpretive assemblage. Agency in that assemblage is genuinely spread, and the spread can be named. The corpus designer fixes what can be found. The model proposes relations within that space. The validating scholar decides whether a proposal is a defensible claim. The participating communities set the criteria by which defensibility is judged for their own traditions. Responsibility does not follow the same lines as agency: it rests with the human validators and with the institutions that deploy the system, since they, not the model, answer for what enters the scholarly record. This asymmetry, distributed agency alongside concentrated responsibility, is the paper's resolution of the tension between calling the system a co-comparatist and retaining human authority over its output.

The harder philosophical question, whether pattern recognition might be turned into interpretation, is not addressed here and perhaps is impossible to decide [13,14]. The pragmatic view taken here is that an idea becomes a defensible

comparative claim when a scholar, employing criteria partly set by the community, determines it has survived scrutiny. That is a lower bar than consciousness and a higher one than mere statistical association, and it locates the threshold in a human judgement rather than in a property of the machine.

4.4 The Four-Layer Interpretive Enhancement Cycle

Governance has to be designed into the research, not bolted on after publication. Figure 4 presents the cycle. Layer 1 is decolonial corpus design, in which communities from marginalised linguistic traditions participate from the outset, not only at a final audit. Layer 2 is autonomous interpretive processing, in which the system generates comparative proposals with stated confidence and traceable evidence. Layer 3 is human scholarly validation, in which interdisciplinary panels accept, refine, or reject proposals. Layer 4 is participatory governance and audit, whose findings revise Layer 1. Table 3 compares the cycle with a standard responsible-AI pipeline and locates the differences: community authorship of the corpus, criteria of defensibility that are tradition-specific rather than universal, and a feedback path that revises the corpus rather than only the model.

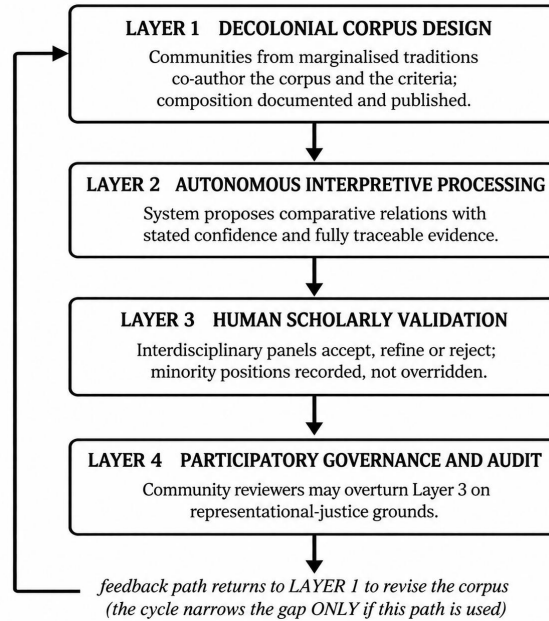


Figure 4. The four-layer interpretive enhancement cycle.

Source: Authors' own elaboration.

Table 3. The proposed cycle compared with a generic human-in-the-loop responsible-AI pipeline.

Attribute	Generic responsible-AI pipeline	Proposed interpretive enhancement cycle
Who authors the data	Engineers and annotators	Marginalised-tradition communities from Layer 1 onward
Criterion of a valid output	Accuracy against a universal ground truth	Defensibility under tradition-specific, community-set standards
Role of the human	Reviews and corrects model output	Judges whether a proposal is a scholarly claim and owns it
Feedback target	Retrains the model	Revises the corpus, then the model
Disagreement	Resolved toward consensus	Preserved as recorded minority positions
Final responsibility	Diffuse or unassigned	Named human validators and deploying institutions

Source: Authors' own elaboration.

The cycle answers a set of operational questions that a schematic version would leave open. Who selects and balances the corpus: mixed teams in which community representatives hold real authority, not advisory seats. How under-representation is determined: through comparison of composition to identified traditions, accompanied by results. How the issue of translation quality and variant texts is addressed: every machine-generated proposal depending on translation is taken as tentative until verified against the source, following the pattern set in cross-cultural validation studies [66]. What information should be presented with a proposal: the texts compared, the model and the version used, the retrieval process, and a statement of confidence. How is acceptability determined: survival of peer review based on community-approved standards. How dissent among validators is managed: minority views are registered, not consensus enforced. Whether Layer 4 can overturn Layer 3: yes, on matters of representational justice, since audit exists

precisely to catch what validation misses. Who bears final responsibility: named human scholars and deploying institutions. How the process is documented: through a retained record of prompts, models, retrieval, and outputs. The claim that each cycle narrows the representational gap is deliberately softened from earlier drafts: the cycle is designed to reduce the gap, and whether it does is an empirical question the design cannot settle in advance.

4.5 A worked literary example

An abstract cycle convinces no one, so consider a concrete pass through it. Suppose a decolonial balanced corpus includes Persian, Malay, and Sicilian narrative collections. In Layer 2 the system proposes that a shared frame-tale structure, a story that contains other stories told to defer a threat, links the three, and it returns the specific passages and a confidence score. In Layer 3 a panel including a Persianist, a Malay literature specialist, and a scholar of medieval Mediterranean texts examines the proposal. Two find it defensible; one objects that the Sicilian case is a later borrowing rather than an independent instance, and the objection is recorded rather than overridden. In Layer 4 a Malay-tradition community notes that the corpus underrepresents oral variants in which the frame works differently, which sends a revision request back to Layer 1. The corpus is rebalanced, the proposal is re-run, and the refined claim, narrower and better evidenced, enters the record with its dissent and its provenance attached. Nothing in this trace requires the machine to understand the tales. It requires the machine to propose, and humans to judge, with communities shaping the ground on which judgement stands. Analogous multi-step, self-directed proposal-and-revision loops are visible in agent-based modelling outside literature [46], offered here only as evidence that such loops are buildable. Figure 5 traces this pass of a proposal through the four layers.

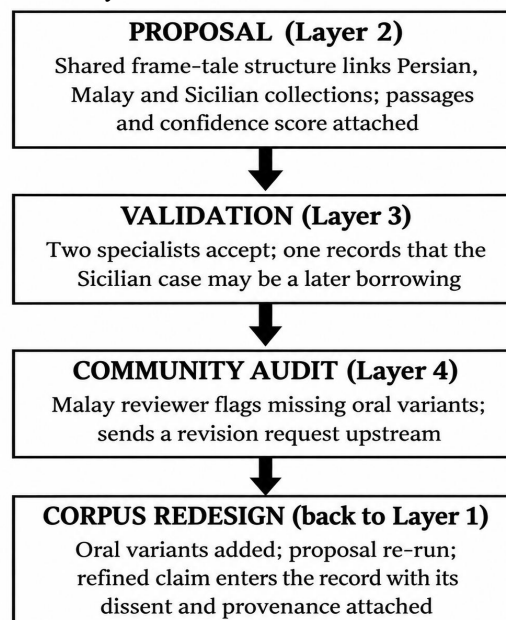


Figure 5. Worked example: a proposal through the cycle.

Source: Authors' own elaboration.

The trace just given assumes everything works, and a reviewer rightly warned that a favourable example proves little. So, consider the same cycle when its conditions fail, and read both passes as illustration rather than as evidence that the framework succeeds. Suppose the corpus is not balanced: the Malay oral variants were never digitised, so Layer 2 never sees them and proposes a frame-tale link that quietly privileges the written Persian and Sicilian forms. Suppose the evidence is not fully traceable, because the model returns a confident relation but cites passages that, on inspection, it has partly fabricated. Suppose no Malay literature specialist is available to the panel, so Layer 3 validates on two traditions and guesses at the third. Suppose the community reviewer at Layer 4 is a single individual presented as representing an entire tradition, and another member of that tradition would have judged differently. And suppose the feedback path is nominal, so the revision request is logged but never funds the digitisation it calls for. Under these conditions the cycle does not merely underperform; it launders a biased proposal through the appearance of due process. This is the failure mode the framework has to be judged against, and it sets real requirements: who counts as a community representative cannot be settled by convenience, and the paper does not pretend it has a general answer; disagreement within a tradition must be recordable rather than resolved by picking one voice; and the whole arrangement is expensive in expertise, time, and access, which means it is feasible only for well-resourced projects unless the cost of multilingual expertise and community participation is met institutionally. Stating these limits is not a concession that weakens the framework. It is what keeps the framework from being the very thing the paper warns against.

4.6 Decolonial Design as More Than Representation

Balancing a corpus by counting texts is the shallowest version of decolonial design, and the paper rejects it as sufficient. The deeper problems are structural. Digitisation is unequal, so some traditions are simply less available to any system [33]. Translation carries politics that shape what a cross-lingual proposal can even see [28]. Intellectual property, community consent, and indigenous data sovereignty govern whether certain materials should enter a corpus at all, questions that epistemic redress models in adjacent domains have begun to formalise [31]. Oral and performance traditions resist textual encoding, so their inclusion cannot be achieved by scanning alone [29]. The cultural assumptions embedded in models and in evaluation criteria are themselves artefacts of skewed training [30,32]. A cycle that treats decolonial design as a numerical target will reproduce the hierarchy it means to dismantle; one that treats it as a governance problem, with communities authoring the corpus and the criteria, has at least a chance of doing otherwise.

4.7 Propositions

The paper's conceptual claims are stated as propositions. Some are falsifiable as written; others depend on constructs that a single study would have to operationalise before it could test them, and those are better read as conceptual expectations than as ready hypotheses. To make the difference usable rather than rhetorical, each proposition below is followed by the construct that would have to be measured and by a sketch of what evidence would support or challenge it. The constructs the reviewer flagged, corpus imbalance, representational-justice objection, validator expertise, proposal quality, and representational gap, are given working definitions here rather than left as intuitions.

P1. The greater a corpus's imbalance toward canonical Western traditions, the more strongly an agentic system's proposals will recommend Western comparators, task held constant.

Measure: corpus imbalance as the ratio of tokens or titles from a named set of Western canonical traditions to the remainder, computed from the documented corpus composition; proposal skew as the share of returned comparators drawn from those same traditions. Support: a positive association across corpora of varying composition. Challenge: skew that persists even when composition is balanced, which would point to model priors rather than corpus.

P2. Proposals whose supporting evidence is fully traceable to source passages survive expert validation at a higher rate than proposals whose evidence is not, independent of novelty.

Measure: traceability as the proportion of a proposal's cited passages that a checker can locate verbatim in the source; validation survival as accept-or-refine rather than reject. Support: higher survival for high-traceability proposals at matched novelty. Challenge: no difference, or novelty dominating traceability.

P3. Proposals that depend on translation for their key comparison are revised or rejected more often than proposals grounded in a shared source language.

Measure: translation dependence as a binary on whether the compared passages share a language; outcome as the validation decision. Support: higher revision or rejection rates for translation-dependent proposals. Challenge: parity, which would ease the translation worry.

P4. Community participation from Layer 1 onward reduces representational-justice objections raised at Layer 4, relative to participation confined to a final audit.

Measure: a representational-justice objection as a logged Layer 4 request citing under-representation, misattribution, or interpretive marginalisation of a tradition; participation timing as an experimental condition. Support: fewer such objections under early participation. Challenge: no reduction, or a rise, which would suggest early participation surfaces rather than prevents objections. This is stated as a conceptual expectation, since running it requires a live deployment.

P5. Acceptance of proposals as defensible claims rises with validator expertise and with the specificity of community-set criteria, not with the system's reported confidence.

Measure: validator expertise as documented standing in the relevant tradition; criteria specificity as the number of explicit, tradition-specific acceptance rules in force; confidence as the system's own score. Support: acceptance tracking expertise and criteria while orthogonal to confidence. Challenge: acceptance tracking the confidence score, which would indicate automation bias.

P6. Distributing agency across the assemblage does not distribute responsibility: perceived accountability stays concentrated on human validators however autonomous the proposal step becomes.

Measure: perceived accountability elicited from practitioners across scenarios that vary the system's autonomy. Support: stable attribution to humans across autonomy levels. Challenge: attribution shifting toward the system as autonomy rises.

P7. Explainability-audited systems yield proposals that validators adjudicate faster than opaque systems, independent of proposal quality.

Measure: adjudication time per proposal; explainability as the presence of an inspectable derivation; quality held constant by matched samples. Support: shorter times for audited systems at equal quality. Challenge: no time difference, which would question the practical value of the audit.

P8. The representational gap narrows only where the Layer 4 to Layer 1 feedback path is actually exercised; where it is nominal, machine participation widens it.

Measure: representational gap as the change in documented corpus composition toward under-represented traditions across cycles; feedback exercise as whether logged Layer 4 requests produced corpus changes. Support: narrowing only in projects that acted on feedback. Challenge: narrowing without exercised feedback, or widening despite it.

Taken together the propositions locate the difference between democratisation and recolonisation not in the technology but in the governance around it, which is the paper's central contention. Figure 6 shows how the eight propositions converge on this outcome.

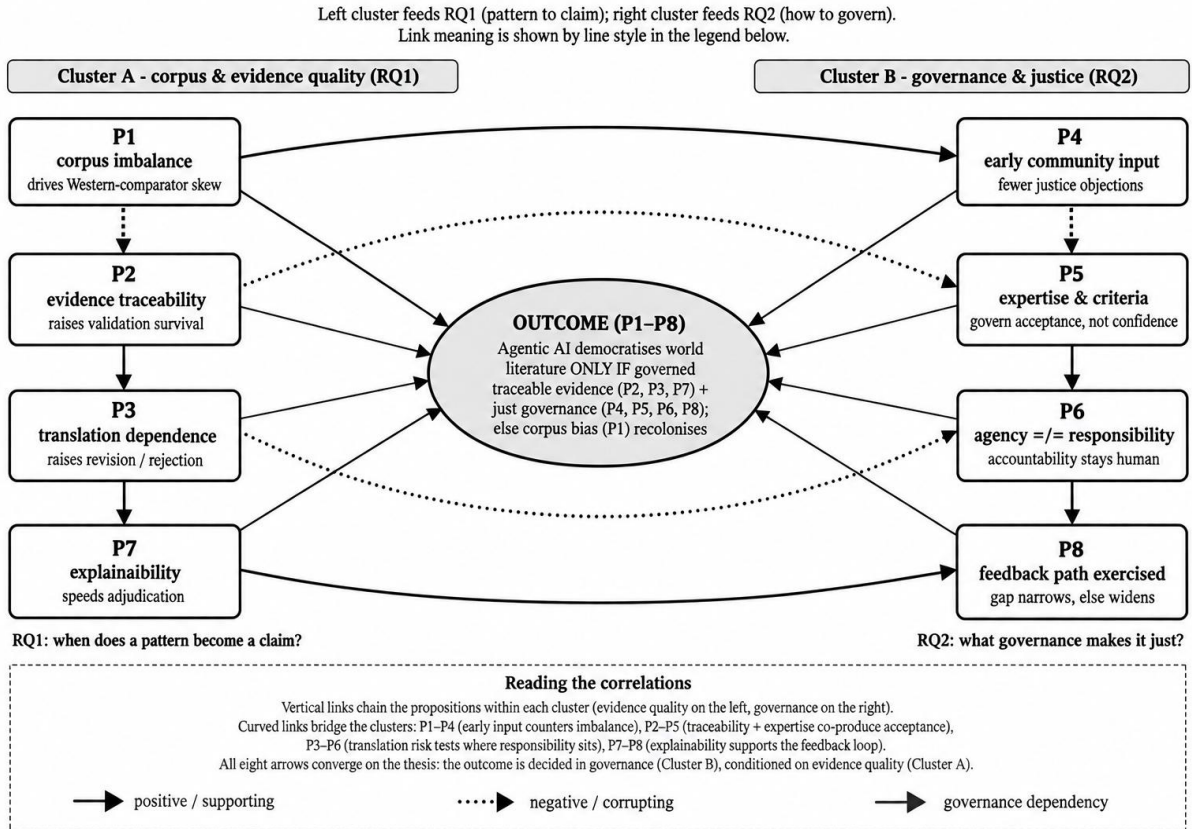


Figure 6. Proposition correlation diagram, showing the outcome derived from all eight propositions.

Source: Authors' own elaboration.

5. Implications, Limitations, and Future Research

5.1 Implications for Theory and Practice

For theory, the paper supplies comparative literature with a specified, rather than gestured-at, account of how a non-human participant could enter interpretation without dissolving human responsibility, and it does so by separating distributed agency from concentrated accountability. It also relocates a familiar worry. The possibility that computational systems may reinforce Eurocentric stratification is already recognized in world literature and computational humanities; the novel contribution of this paper is not an announcement of this possibility, but rather the anchoring of this possibility in concrete and reviewable design choices. Adoption research across domains reveals just how quickly this happens when the system proves useful [18,64,65,75], and the same is true for AI adoption in hospitality, retail, and services [76-86]. Research on attitudes towards AI in emotions and decision-making highlights that acceptance is not consistent across society and demographics [87]; that decision support through AI alters judgement in small organizations [88], and by extension, that reliance on advisory AI tools like robo advisors influences the way that principals evaluate consequences [89]. The practical takeaway for the research community is limited but clear-cut: release the corpus makeup, keep the source of each proposal, and empower the community review process.

5.2 Limitations

The framework has not been tested, and its weaknesses are of two kinds that the paper now keeps apart, since a reviewer rightly noted that some of these threats strike at the framework's credibility rather than sitting beside it. The first kind can be reduced by the governance the paper proposes, though not abolished. Hallucinated quotations and attributions, which are disqualifying in a field where the exact words carry the argument, are the reason Layer 3 requires every cited passage to be checked verbatim against the source before a proposal is accepted; the check catches fabrication but adds cost and depends on the checker's diligence. Translation error is handled the same way, by treating any translation-dependent proposal as provisional until verified against the source [2], which lowers but does not remove the risk. Automation bias, the tendency to defer to a confident machine, is the reason acceptance criteria are tied to validator expertise and community rules rather than to the system's confidence score, and P5 is designed to detect it if it persists.

Opaque licensing and copyright constraints on what may enter Layer 1 are a governance matter of documentation and permission rather than a technical limit. The second kind cannot be resolved by governance and bounds the whole proposal. Prompt sensitivity and model instability mean the same request can yield different proposals on different days [45], and version dependence makes exact reproducibility fragile [43]; governance can log these but cannot make an inherently stochastic system deterministic. Explainability is only partial, so validators may adjudicate outputs whose derivation they cannot fully inspect [44], and no auditing regime closes that gap entirely. Deeper still, alignment and epistemic-limit arguments hold that some of what interpretation needs may lie outside what current systems can supply at all [14,47,48], and the uncanny-valley evidence suggests, by analogy only, that trust in machine-made cultural output is itself unstable [90].

The reliability of the surrounding literature is uneven as well; some adoption-side work, such as data-science approaches to encouraging entrepreneurship [91], is available only as a preprint that has not yet been peer reviewed, a reminder that inclusion in a corpus does not by itself certify reliability. Governance experience from adjacent high-stakes deployments, including AI-driven healthcare chatbots [92] and autonomous-service business ecosystems [93], indicates that documented provenance and named human sign-off are the floor any credible arrangement requires, and that pedagogies building ethical judgement through repeated practice suggest how validators might be trained rather than assumed [94]. The honest summary is that the first set of limitations is manageable and the second is not, and that the framework is worth pursuing only where the manageable ones are actually managed.

5.3 Future Research

A direct follow-up step would involve a validation study involving the comparison of machine-generated comparative proposals against human expertise across multiple traditions, using existing tools from cross-cultural design research [66] and evaluating P2, P3, and P5. Further research could involve longitudinal analysis of whether standards of validity change within disciplines due to machine involvement [18]. Comparative research across different academic communities would involve investigating the reception of machine proposals within other traditions, as there are known differences in attitudes toward AI [32,80]. Governance research should test the participatory model empirically rather than asserting it, and pedagogical research should ask how comparatists are trained for this role, extending education-focused studies of AI-supported assessment, validity, and fairness [34,51,70,95-97] and of AI-mediated entrepreneurship and creativity in learning [75,97]. Design-oriented studies of generative systems in visual and spatial domains [45,98-100] and reviews of AI in high-stakes technical settings [44,58,101-103] may supply, by analogy only, method for auditing self-directed systems. The accounting-style question of how interpretive credit is attributed across a distributed assemblage [54,104] becomes sharper, not softer, as more of the work is delegated.

6. Conclusion

This paper began from an anomaly. Comparative literature defines itself by border-crossing yet has no account of the border between human and machine interpretation, precisely as that border is being pushed. The response offered here is neither to celebrate nor to dismiss, but to specify. An agentic system can occupy the functional role of a co-comparatist, originating comparative proposals that humans must judge, and that functional shift is enough to unsettle anthropocentric method without granting the machine understanding, consciousness, or responsibility.

The argument has been kept deliberately conditional. The question of whether machine involvement enlarges world literature or recolonizes it does not rest with the technology itself, which is why the paper will not take a stance regarding how agentic AI already transforms the discipline. Rather, the paper argues that agentic AI forces the researcher to make a choice, and that the choice exists in the domain of governance in whose hand is the authoring of the corpus, in what constitutes a valid claim, and in whether there is indeed a feedback loop from audit to design. The reverse outcome is given equal standing throughout, for the reason that a machine learning on the texts available right now and operating without the feedback loop outlined above would simply reinforce the canon in a more powerful way. The window within which these architectural choices are possible is not limitless, but neither is it quite as narrow as fear would imply, and the only realistic assessment is that what lies ahead for the field is constructive work.

Acknowledgements

The authors thank the editor and the reviewers, whose detailed and demanding comments materially improved the conceptual precision, disciplinary grounding, and citation logic of this paper.

Funding

This research received no external funding.

Ethics Statement

Not applicable. This is a conceptual study based solely on published literature; it did not involve human participants, human data, or animals, and therefore did not require approval from an ethics committee or institutional review board.

Informed Consent

Not applicable. The study did not involve human participants.

Data Availability Statement

No new primary data were created or analysed in this study. All sources analysed are published works listed in the reference list. The search records and screening documentation supporting the corpus are available upon reasonable request from the corresponding author.

Author Contributions

All authors shared the conceptualisation, the methodology, and the development of the framework. The first author led the drafting of the introduction, the literature review, the theory development, and the worked example. The second author led the conceptual method, the evidential audit, the propositions, and the limitations. Both authors revised the manuscript in response to review, and read and approved the final version.

Conflict of Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Generative AI Statement

Generative AI tools were used to assist the authors during the preparation of this manuscript. Specifically, a large language model was used to support language editing, to help structure the argument, and to assist with formatting the reference list. All conceptual claims, the framework, the propositions, the worked example, and the selection and interpretation of sources are the authors' own. The authors reviewed, verified, and take full responsibility for the entire content, including every citation and its relationship to the claim it supports.

References

- [1] Xin Y, Begaliyeva SB. Influence of short novels on creation of educational programs in literature: Taking A. P. Chekhov's 'The Chameleon' and Lu Xun's 'A Madman's Diary' as examples. *Education Sciences*, 2025, 15(7), 906. DOI: 10.3390/educsci15070906
- [2] Mohamed YA, Khanan A, Bashir M, Mohamed AHM, Adiel MAE, Elsadig MA. The impact of artificial intelligence on language translation: A review. *IEEE Access*, 2024, 12, 25553-25579. DOI: 10.1109/ACCESS.2024.3366802
- [3] Xia Y, Shin SY, Kim JC. Cross-cultural intelligent language learning system (CILS): Leveraging AI to facilitate language learning strategies in cross-cultural communication. *Applied Sciences*, 2024, 14(13), 5651. DOI: 10.3390/app14135651
- [4] Burriss SK, Leander K. Critical posthumanist literacy: Building theory for reading, writing, and living ethically with everyday artificial intelligence. *Reading Research Quarterly*, 2024, 59(4), 560-569. DOI: 10.1002/rrq.565
- [5] Wang Z, Wang C. A posthumanist approach to AI literacy. *Computers and Composition*, 2025, 76, 102933. DOI: 10.1016/j.compcom.2025.102933
- [6] Lotherington H, Pegrum M, Thumlert K, Tomin B, Boreland T, Pobuda T. Exploring opportunities for language immersion in the posthuman spectrum: Lessons learned from digital agents. *Interactive Technology and Smart Education*, 2025, 22(4), 525-547. DOI: 10.1108/ITSE-02-2024-0038
- [7] Hosseini S, Seilani H. The role of agentic AI in shaping a smart future: A systematic review. *Array*, 2025, 26(1), 100399. DOI: 10.1016/j.array.2025.100399
- [8] Tirulo A, Yadav M, Lolamo M, Chauhan S, Siano P, Shafie-khah M. Beyond automation: Unveiling the potential of agentic intelligence. *Renewable and Sustainable Energy Reviews*, 2025, 226. DOI: 10.1016/j.rser.2025.116218

- [9] Holldack F, Banh L, Strobel G. Agentic information systems. *Electronic Markets*, 2026, 36, 5. DOI: 10.1007/s12525-025-00861-0
- [10] Raghavendra P, Saikia MJ. Agentic AI: A perspective on architecture, frameworks and applications. *AI*, 2026, 7(6), 219. DOI: 10.3390/ai7060219
- [11] Torres-Martinez S. Credence, attribution, and creativity in the construction of literary meaning with generative artificial intelligence. *Language and Semiotic Studies*, 2026, 12(2), 240-280. DOI: 10.1515/lass-2024-0072
- [12] Kumar PC, Cotter K, Cabrera LY. Taking responsibility for meaning and mattering: An agential realist approach to generative AI and literacy. *Reading Research Quarterly*, 2024, 59(4), 570-578. DOI: 10.1002/rrq.570
- [13] Farina M, Wang Y, Kladko S. Ethical and epistemological reflections on autonomous AI-powered agents (AAIAs). *Topoi*, 2025, 1-17. DOI: 10.1007/s11245-025-10218-z
- [14] Stranak P. What artificial intelligence may be missing and why it is unlikely to attain it under current paradigms. *Philosophies*, 2026, 11(1), 20. DOI: 10.3390/philosophies11010020
- [15] Spennemann DHR. Generative artificial intelligence, human agency and the future of cultural heritage. *Heritage*, 2024, 7(7), 3597-3609. DOI: 10.3390/heritage7070170
- [16] Yang J, Liu T, Luo YT, Pang PCI. Deep learning in cultural imagery dissemination: A systematic scoping review of AI-driven visual transmission mechanisms. *Frontiers in Communication*, 2025, 10, 1645168. DOI: 10.3389/fcomm.2025.1645168
- [17] Wiboolyasarin W, Wiboolyasarin K, Tiranant P, Jinowat N, Boonyakitanont P. AI-driven chatbots in second language education: A systematic review of their efficacy and pedagogical implications. *Ampersand*, 2025, 14, 100224. DOI: 10.1016/j.amper.2025.100224
- [18] Riandhi AN, Arviansyah MR, Sondari MC. AI and consumer behavior: Trends, technologies, and future directions from a Scopus-based systematic review. *Cogent Business & Management*, 2025, 12(1). DOI: 10.1080/23311975.2025.2544984
- [19] Hayles NK. *How we became posthuman: Virtual bodies in cybernetics, literature, and informatics*. Chicago: University of Chicago Press, 1999.
- [20] Braidotti R. *The Posthuman*. Cambridge, UK: Polity Press, 2013.
- [21] Latour B. *Reassembling the social: An introduction to actor-network-theory*. Oxford University Press, 2005.
- [22] Iapaolo F, Lynch CR. Rethinking urban AI: A critical posthumanist perspective on intelligence, autonomy, and agency. *Urban Geography*, 2026, 47(1), 65-87. DOI: 10.1080/02723638.2025.2510812
- [23] Halpin H. Artificial intelligence versus collective intelligence. *AI & Society*, 2025, 40, 4589-4604. DOI: 10.1007/s00146-025-02240-x
- [24] Damrosch D. *What is world literature?* Princeton University Press, 2003.
- [25] Moretti F. *Distant reading*. Verso, 2013.
- [26] Casanova P. *The world republic of letters*. Harvard University Press, 2004.
- [27] Apter E. *Against world literature: On the politics of untranslatability*. Verso, 2013.
- [28] Mao L, Kong B, Chu S. Intercultural communicative competence in translation teaching mediated by artificial intelligence: Cultural schema theory, mediated discourse analysis and intercultural pragmatics. *Acta Psychologica*, 2026, 265, 106616. DOI: 10.1016/j.actpsy.2026.106616
- [29] Suryanto TLM, Wibawa AP, Hariyono, Nafalski A, Kurubacak Cakir G. Generated cultural heritage question-answer dataset: Durga in multi-dimensional perspectives. *Data in Brief*, 2026, 65. DOI: 10.1016/j.dib.2026.112495
- [30] Deng M, Chen L, CDGFD: Cross-domain generalization in ethnic fashion design using LLMs and GANs: A symbolic and geometric approach. *IEEE Access*, 2025, 13, 7192-7207. DOI: 10.1109/ACCESS.2024.3524444
- [31] Maqutu L. Epistemic decolonisation and the integration of African traditional medicine: Toward an epistemic redress model for inclusive AI in mental healthcare. *AI & Society*, 2026. DOI: 10.1007/s00146-026-03182-8
- [32] Xue Y, Chinapah V, Zhu C. A comparative analysis of AI privacy concerns in higher education: News coverage in China and Western countries. *Education Sciences*, 2025, 15(6), 650. DOI: 10.3390/educsci15060650
- [33] Yao X, Lin G. Reframing digital literacy: Current paradigms, core challenges, and future research directions. *Humanities & Social Sciences Communications*, 2026, 13, 1130. DOI: 10.1057/s41599-026-07496-2
- [34] Nopas DS. Algorithmic learning or learner autonomy? Rethinking AI's role in digital education. *Qualitative Research Journal*, 2025. DOI: 10.1108/QRJ-11-2024-0282
- [35] Bode K. *A world of fiction: Digital collections and the future of literary history*. University of Michigan Press, 2018.
- [36] Underwood T. *Distant horizons: Digital evidence and literary change*. University of Chicago Press, 2019.
- [37] Risam R. *New digital worlds: Postcolonial digital humanities in theory, praxis, and pedagogy*. Northwestern University Press, 2019.
- [38] Said EW. *Orientalism*. Pantheon Books, 1978.
- [39] Spivak GC. Can the subaltern speak? *Marxism and the Interpretation of Culture*, 1988, 271-313.
- [40] Venuti L. *The translator's invisibility: A history of translation*. Routledge, 1995.
- [41] Bassnett S. *Translation studies*. Routledge, 1980.
- [42] Ahmadzadeh S, Karthick G, Alsafi T. Large language models for future internet ecosystems: A taxonomy-based review of smart IoT, Edge intelligence, and autonomous AI. *Future Internet*, 2026, 18(8), 393. DOI: 10.3390/fi18080393
- [43] Jing F, Zhang Y, Gao M, Zhang X, Zhou H. A review of federated large language models for Industry 4.0. *Sensors*, 2026, 26(4), 1116. DOI: 10.3390/s26041116
- [44] Yadav A, Srivastava V, Yadav A. Guided relevance attention mapping: Explainable artificial intelligence reimaged. *Engineering Applications of Artificial Intelligence*, 2025, 163. DOI: 10.1016/j.engappai.2025.112925
- [45] Colombo G, Niederer S, Gaetano CD. From prompt engineering to prompt design: Research strategies for visual generative AI. *Big Data & Society*, 2026, 13(2), 1-19. DOI: 10.1177/20539517261451462
- [46] Vannier C, Dete M, Richards D, Heays A, Lavorel S, Quenol H. From rules to roles: A generative AI-driven agent-based model for land-use decisions. *Ecological Modelling*, 2026, 521. DOI: 10.1016/j.ecolmodel.2026.111731
- [47] Bruiger D. Reflections on the AI alignment problem. *AI & Society*, 2025, 40(6), 4383-4392. DOI: 10.1007/s00146-025-02211-2
- [48] Ali F, Bouzoubaa K, Gelli F, Hamzi B, Khan S. Islamic ethics and AI: An evaluation of existing approaches to AI using trusteeship ethics. *Philosophy & Technology*, 2025, 38, 120. DOI: 10.1007/s13347-025-00922-4

- [49] Rezaei Z, Khorraminia A, Shi D, Banad YM. Network-based artificial intelligence in mental healthcare: A systematic review of chatbots, artificial intelligence/machine learning models and ethical considerations in global healthcare networks. *Digital Health*, 2026, 12. DOI: 10.1177/20552076261421688
- [50] Kasubi JW, Kisumbe LA, Nyabakora WI. Mapping the knowledge base for the impact of artificial intelligence on human resources management: A bibliometric study. *SAGE Open*, 2025, 15(3). DOI: 10.1177/21582440251377298
- [51] Helal MYI, Elgendy IA, Albashrawi MA, Dwivedi YK, Al-Ahmadi MS, Jeon I. The impact of generative AI on critical thinking skills: A systematic review, conceptual framework and future research directions. *Information Discovery and Delivery*, 2025. DOI: 10.1108/IDD-05-2025-0125
- [52] Helal MY, Elgendy IA, Albashrawi MA, Chae I, Nusair K, Dwivedi YK. Mapping current artificial intelligence research trends in hospitality: A comprehensive bibliometric review. *Information Discovery and Delivery*, 2026, 1-17. DOI: 10.1108/IDD-09-2025-0244
- [53] Nichifor B, Zait L, Timiras L. Drivers, barriers, and innovations in sustainable food consumption: A systematic literature review. *Sustainability*, 2025, 17(5), 2233. DOI: 10.3390/su17052233
- [54] Sun Y, Gou D. Trust, trustworthiness and credibility in corporate social media communication: A TCM-ADO framework-based systematic review. *Journal of Business & Industrial Marketing*, 2026. DOI: 10.1108/JBIM-05-2025-0477
- [55] Utami RD, Aimin W. Balancing personalization, privacy, and value: A systematic literature review of AI-enabled customer experience management. *Information*, 2026, 17(2), 115. DOI: 10.3390/info17020115
- [56] Hassan M, Shrabani SS, Islam MA, Basaruddin KS, Ijaz MF, Bin Kamarrudin NS, et al. Integration of extended reality technologies in transportation systems: A bibliometric analysis and review of emerging trends, challenges, and future research. *Results in Engineering*, 2025, 26, 105334. DOI: 10.1016/j.rineng.2025.105334
- [57] Ahn KU, Kim DW, Cho HM, Chae CU. Alternative approaches to HVAC control of chat generative pre-trained transformer (ChatGPT) for autonomous building system operations. *Buildings*, 2023, 13(11), 2680. DOI: 10.3390/buildings13112680
- [58] Saeed RS, Raja AA, Iqbal R, Nizami IF, Shahzad MS, Yaseen M. Adaptive and intelligent control architectures for smart microgrid systems: A comprehensive review. *Energy Conversion and Management: X*, 2026, 30, 101772. DOI: 10.1016/j.ecmx.2026.101772
- [59] Lazarowska A. A comparative analysis of computational intelligence methods for autonomous navigation of smart ships. *Electronics*, 2024, 13(7), 1370. DOI: 10.3390/electronics13071370
- [60] Mecca B, Lami IM, Cugurullo F. Smart, autonomous and sustainable cities? A critical analysis from the perspective of strong, weak and super weak sustainability. *Cities*, 2026, 174. DOI: 10.1016/j.cities.2026.107099
- [61] Li J, Li J, Yang G, Yang L, Chi H, Yang L. Applications of large language models and multimodal large models in autonomous driving: A comprehensive review. *Drones*, 2025, 9(4), 238. DOI: 10.3390/drones9040238
- [62] Khatab E, Fathy F, AlKholy A, Shalash O. Sensor fusion and perception for autonomous driving: A critical review of modalities, AI models, algorithms, and industry configurations. *Machine Learning and Knowledge Extraction*, 2026, 8(7), 199. DOI: 10.3390/make8070199
- [63] Viktor P, Kiss G. Sensors in self-driving vehicles: A detailed literature review and new trends. *Sensors*, 2026, 26(7), 2153. DOI: 10.3390/s26072153
- [64] Na S, Heo S, Choi W, Kim C, Whang SW. Artificial intelligence (AI)-based technology adoption in the construction industry: A cross-national perspective using the technology acceptance model. *Buildings*, 2023, 13(10), 2518. DOI: 10.3390/buildings13102518
- [65] Sathasivam K, Tan MK, Thomas A, Koon VY. Unveiling human resource's key role in driving organizational digital transformation: A qualitative inquiry. *European Journal of Innovation Management*, 2026, 29(11), 295-316. DOI: 10.1108/EJIM-08-2025-1065
- [66] Rodriguez-Rodriguez AM, Fuente-Costa MD, de la Riva ME, Perez-Dominguez B, Hernandez-Sanchez S, Paseiro-Ares G, et al. Spanish language version of the 'medical quality video evaluation tool' (MQ-VET): Cross-cultural AI-supported adaptation and validation study. *Science Progress*, 2025, 108(1), 368504251327507. DOI: 10.1177/00368504251327507
- [67] Feng S, Ma S, Zhu X, Yan M. Artificial intelligence-based underwater acoustic target recognition: A survey. *Remote Sensing*, 2024, 16(17), 3333. DOI: 10.3390/rs16173333
- [68] Ming Y, Yang XR, Wang WH, Chen Z, Feng JL, Xing YF, et al. Benchmarking neural radiance fields for autonomous robots: An overview. *Engineering Applications of Artificial Intelligence*, 2024, 140, 109685. DOI: 10.1016/j.engappai.2024.109685
- [69] Bodige R, Akarapu R, Poladi PK. Transformer-based advances in sarcasm detection: A study of contextual models and methodologies. *Knowledge and Information Systems*, 2025, 67(9), 7399-7430. DOI: 10.1007/s10115-025-02469-4
- [70] Sun J, Rui J. Evaluating the longitudinal effects of AI-enhanced collaborative dialogue modes on computational thinking and language proficiency in EFL learners: A mixed-methods approach. *Journal of Educational Computing Research*, 2026, 64(3), 539-570. DOI: 10.1177/07356331251399452
- [71] Hassan B, Raza MO, Siddiqi Y, Wasiq MF, Siddiqui RA. CONNECT: An AI-powered solution for student authentication and engagement in cross-cultural digital learning environments. *Computers*, 2025, 14(3), 77. DOI: 10.3390/computers14030077
- [72] Cheng CY, Hsu CH. Substitution or complementarity? Understanding the role of ChatGPT in transforming news media traffic in the United States and Taiwan. *Data Technologies and Applications*, 2025, 59(4-5), 641-669. DOI: 10.1108/DTA-02-2025-0151
- [73] Pang Y, Ma Q, Hasan N. Utilizing big data and artificial intelligence to improve the cross-border trade English education. *PLOS One*, 2025, 20(11), e0323941. DOI: 10.1371/journal.pone.0323941
- [74] Morden JN, Caraffini F, Kypraios I, Al-Bayatti AH, Smith R, Tan YA. Driving in the rain: A survey toward visibility estimation through windshields. *International Journal of Intelligent Systems*, 2023(1), 9939174. DOI: 10.1155/2023/9939174
- [75] Ganuthula VRR. AI-enabled individual entrepreneurship theory: Redefining scale, capability, and sustainability in the digital age. *Journal of Innovation and Entrepreneurship*, 2025, 14, 85. DOI: 10.1186/s13731-025-00521-9
- [76] Xing Y, Zhang JZ. The algorithmic guest: AI as a co-creator in customer experience management. *International Journal of Contemporary Hospitality Management*, 2026, 38(4), 1433-1452. DOI: 10.1108/IJCHM-07-2025-1138
- [77] Jaiswal R, Gupta S, Chen PJ. GenAI personalization: Antecedents, outcomes, mediators, and moderators. *International Journal of Contemporary Hospitality Management*, 2026, 38(13), 134-158. DOI: 10.1108/IJCHM-07-2025-1056

- [78] Song C, Zhou S. Push or pull? Understanding switching intentions from human services to GAI agents through the PPM framework. *Computers in Human Behavior*, 2025, 176. DOI: 10.1016/j.chb.2025.108874
- [79] Campbell C, Sands S, Mavrommatis A. The paradoxes of generative AI-mediated acculturation. *International Marketing Review*, 2026. DOI: 10.1108/IMR-10-2025-0569
- [80] Chow M, Ho S, Armutcu B, Tan A, Butt MK. Role of culture in how AI affects the brand experience: Comparison of Belt and Road countries. *Asia-Pacific Journal of Business Administration*, 2025, 18(1), 213-236. DOI: 10.1108/APJBA-09-2024-0510
- [81] Wang YC, Zhang YT, Ma E, Zhang L, Zhang Y, Ning H, et al. AI-assisted cross-cultural training for hotel newcomers: An IT mindfulness-driven perspective on building cultural intelligence and career development confidence. *International Journal of Hospitality Management*, 2025, 128, 104186. DOI: 10.1016/j.ijhm.2025.104186
- [82] Tran MT. Advancing retail and service strategies: AI-driven consumer behavior prediction, gamification, and ethical marketing. *Journal of Retailing and Consumer Services*, 2025, 88. DOI: 10.1016/j.jretconser.2025.104558
- [83] Miskolczi M, Michalko GI, Farkas J. Travel into my future: Perceptions of Generation Z along a fictive journey. *Journal of Tourism Futures*, 2025, 1-22. DOI: 10.1108/JTF-07-2024-0141
- [84] Shin H, Xie R, Yoon H, Lee J. Measuring generative AI-mediated travel experiences: Development and validation of a multidimensional scale. *Tourism Management Perspectives*, 2026, 63. DOI: 10.1016/j.tmp.2026.101492
- [85] Manzoor MF, Afraz MT, Waseem M, Ahmed Z. Psychology of eating the future: Consumer acceptance, digital influence and behavioral drivers of novel foods. *Foods*, 2026, 15(14), 2471. DOI: 10.3390/foods15142471
- [86] Feng X, Lee K, Madanoglu M. Exploring Pinot Noir consumption among younger Chinese consumers. *International Journal of Wine Business Research*, 2026, 1-27. DOI: 10.1108/IJWBR-01-2026-0009
- [87] Mantello P, Ho MT, Nguyen M, Vuong QH. Bosses without a heart: Socio-demographic and cross-cultural determinants of attitude toward emotional AI in the workplace. *AI & Society*, 2023, 38(1), 97-119. DOI: 10.1007/s00146-021-01290-1
- [88] Kugara G. AI-augmented decision-making and sustainability performance in SMEs: A hybrid PLS-SEM, ANN, and multi-group analysis. *Industrial Management & Data Systems*, 2026, 1-24. DOI: 10.1108/IMDS-11-2025-1719
- [89] Bihari A, Dash M, Kumar A, Muduli K, Luthra S, Samadhiya A. The impact of AI-powered robo-advisors on investor decision-making: The moderating role of financial knowledge. *VINE Journal of Information and Knowledge Management Systems*, 2026, 1-30. DOI: 10.1108/VJIKMS-09-2025-0407
- [90] Scaffidi Abbate C, Taddeo L, Di Nuovo S. Humanoid interfaces in artificial intelligence-based language learning devices: Possible 'uncanny valley' effects? *Acta Psychologica*, 2025, 256, 104997. DOI: 10.1016/j.actpsy.2025.104997
- [91] Chanthathi SR. A segmented approach to encouragement of entrepreneurship using data science. *Authorea Preprints*, 2024. DOI: 10.22541/au.172115314.43206877/v1
- [92] Wah JNK. Revolutionizing e-health: The transformative role of AI-powered hybrid chatbots in healthcare solutions. *Frontiers in Public Health*, 2025, 13, 1530799. DOI: 10.3389/fpubh.2025.1530799
- [93] Steiber A, Alvarez D. AI-driven digital business ecosystems: A study of Haier's EMCs. *European Journal of Innovation Management*, 2024, 28(8), 3966-3984. DOI: 10.1108/EJIM-01-2024-0076
- [94] Ekasari K, Djamhuri A, Eltivia N, Miharso A. Ethics in motion: A transformative pedagogy of repetition for future-ready accounting education. *International Journal of Ethics and Systems*, 2026, 1-34. DOI: 10.1108/IJOES-07-2025-0404
- [95] Zacharis G, Papadakis S. Can AI grade like a human? Validity, reliability, and fairness in university coursework assessment. *Educational Process: International Journal*, 2025, 19, e2025591. DOI: 10.22521/edupij.2025.19.591
- [96] Ugras H, Ugras M, Papadakis S, Kalogiannakis M. ChatGPT-supported education in primary schools: The potential of ChatGPT for sustainable practices. *Sustainability*, 2024, 16(22), 9855. DOI: 10.3390/su162209855
- [97] Jeong-Hyun P, Seon-Joo K, Sung-Tae L. AI and creativity in entrepreneurship education: A systematic review of LLM applications. *AI*, 2025, 6(5), 100. DOI: 10.3390/ai6050100
- [98] Lou Y, Norman D, Thackara J, Sotamaa Y, Somerson R, Imbesi L, et al. DesignS: World design cities Shanghai Manifesto 2024. *She Ji: The Journal of Design, Economics, and Innovation*, 2025, 11(2), 236-246. DOI: 10.1016/j.sheji.2025.05.003
- [99] Liu S, Liu D. The application of AI-driven color theory in graphic design and branding. *International Journal of Computational Intelligence Systems*, 2026, 19, 229. DOI: 10.1007/s44196-026-01298-9
- [100] Zhao CW, Yang J, Li J. Generation of hospital emergency department layouts based on generative adversarial networks. *Journal of Building Engineering*, 2021, 43, 102539. DOI: 10.1016/j.jobbe.2021.102539
- [101] Amarasooriya RP, Gregory MA, Li S. A complete survey on artificial intelligence-based resource allocation for sixth generation mobile. *Ad Hoc Networks*, 2026, 189, 104264. DOI: 10.1016/j.adhoc.2026.104264
- [102] Farhat H, Altarawneh A. Physics-informed machine learning for intelligent gas turbine digital twins: A review. *Energies*, 2025, 18(20), 5523. DOI: 10.3390/en18205523
- [103] Eti E, Kardes O, Yuksel S, Dincer H, Ertosun GO, Doguc O. An artificial intelligence-driven fuzzy decision-making framework integrating natural language processing and rule-based logic for preventing agile fatigue. *Artificial Intelligence Review*, 2026, 59, 163. DOI: 10.1007/s10462-026-11562-1
- [104] Steingard D, Linacre S, Reibstein D, Flynn D, Springer C. Revolutionizing societal impact in business school research: Can the FT50 lead the change? *Society and Business Review*, 2025, 21(2), 316-333. DOI: 10.1108/SBR-08-2025-0334