



Castledown

 OPEN ACCESS

Technology in Language Teaching & Learning

ISSN 2652-1687

<https://www.castledown.com/journals/tltl/>

Technology in Language Teaching & Learning, 7(4), 103824 (2025)

<https://doi.org/10.29140/tltl.v7n4.103824>

Framing Human–AI Collaboration in Language Education



GLENN STOCKWELL^a 

YIJEN WANG^b 

^a*The Education University of Hong Kong, Hong Kong*
gstock@eduhk.hk

^b*Waseda University, Japan*
y.wang@aoni.waseda.jp

Abstract

The increasing use of generative artificial intelligence in language education has been accompanied by an increasing tendency to describe human-AI interaction as “collaboration.” This paper examines that framing by first clarifying what collaboration entails and then considering the extent to which human-AI interaction aligns with these requirements. Drawing on work in collaborative learning, human-AI teaming, and workplace collaboration, the paper identifies seven commonly cited features of collaboration: shared goals and purpose, interdependence and coordination, reciprocal communication, distributed accountability, complementary expertise, calibrated trust, and acknowledged contribution. While human-AI interaction in language education displays surface similarities to collaboration, particularly in terms of task interdependence and functional complementarity, it differs in more fundamental respects. In particular, human-AI interaction lacks shared purpose, reciprocal understanding, and distributed accountability, all of which are central to established conceptions of collaboration. These distinctions matter for how teachers, learners, and researchers interpret AI use and for how responsibility, evaluation, and learning are understood in AI-mediated activity. The paper argues that greater conceptual precision around collaboration is necessary if claims about human-AI relationships in language education are to remain theoretically grounded and pedagogically meaningful.

Keywords: Generative AI, collaboration, language education, human-AI interaction, teacher education

Copyright: © 2025 Glenn Stockwell & Yijen Wang. This is an open access article distributed under the terms of the Creative Commons Attribution Non-Commercial 4.0 International License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.
Data Availability Statement: All relevant data are within this paper.

Introduction

The emergence of generative artificial intelligence tools, particularly large language models such as ChatGPT, has prompted a striking shift in how teachers, learners, and researchers describe their engagement with educational technology. Where earlier discussions emphasised tools, resources, and assistive functions, current discourse increasingly frames human-AI interaction as “collaboration”. Teachers describe collaborating with AI to develop materials; learners report collaborating with chatbots on writing and practice tasks; and researchers acknowledge AI as a collaborator in manuscript preparation. This buzzword now appears not only in everyday educational practice, but also in institutional guidelines and academic publications.

The use of collaboration language to describe human–machine interaction is not new. Research in human–computer interaction has long explored the idea of human–computer collaboration, particularly in contexts where computational systems appear to participate actively in problem solving. Silverman’s (1992) early work examined collaboration as a design aspiration in human–computer systems, focusing on coordination, role allocation, and decision support. Crucially, however, this work did not treat human–computer collaboration as equivalent to human–human collaboration. Instead, collaboration was framed as a functional relationship shaped by system design and task structure, rather than by shared understanding, intention, or responsibility.

Subsequent Human-Computer Interaction (HCI) research has consistently highlighted the asymmetries inherent in human–machine interaction. Analyses grounded in distributed cognition emphasise that apparent collaboration with computational systems often reflects the distribution of cognitive work across people, artefacts, and environments, rather than genuinely shared agency or meaning-making (Hollan et al., 2000). Related work on situated action cautions against attributing intention or understanding to machine behaviour, however adaptive it may appear, arguing that such attributions risk obscuring the fundamentally different bases of human and machine action (Suchman, 1987, 2007).

More recent work in human–AI interaction continues this cautious stance. Contemporary HCI research emphasises the importance of maintaining human oversight, responsibility, and control in AI-supported activity, even where systems appear highly interactive or autonomous (Amershi et al., 2019; Shneiderman, 2020). Rather than treating AI systems as collaborators in a strong sense, this literature frames effective human–AI engagement in terms of calibrated interaction, appropriate reliance, and clearly bounded roles. Even in domains where human and computational contributions are tightly coupled, collaboration is typically understood as asymmetrical and structured, rather than reciprocal or jointly owned (Kittur et al., 2013).

The renewed use of collaboration language in relation to generative AI therefore raises important questions. Does the term accurately capture what occurs when humans interact with generative AI? Or does it obscure distinctions that matter for how AI use is understood, evaluated, and taught? These questions are particularly pressing in language education, where learners are expected to develop communicative competence through interaction and effort, yet where AI tools can produce fluent language output with minimal learner engagement. How collaboration is framed therefore has implications not only for teachers’ and researchers’ practices, but also for how learners interpret their own roles, responsibilities, and learning processes when working with AI.

This paper examines whether “collaboration” is an appropriate descriptor for human-AI interaction in language education. Rather than assuming an answer, it begins by articulating what collaboration requires, drawing on established theoretical frameworks. It then considers human-AI interaction in relation to these requirements from the perspectives of language teachers, learners, and researchers.

The paper suggests that while some features commonly associated with collaboration are present in human-AI interaction, others are fundamentally absent. These distinctions are not merely semantic; they shape how teachers design learning tasks, how learners engage with AI in ways that support or undermine development, and how researchers evaluate claims about AI-enhanced learning.

The paper focuses specifically on generative AI tools, specifically those capable of producing text, code, or other content in response to prompts, rather than the broader category of AI in education that includes intelligent tutoring systems, automated assessment, and adaptive learning platforms. This focus reflects both the recency of generative AI's availability and the particular ways in which it has encouraged collaboration language in educational discourse.

Defining Collaboration

Collaboration is not simply working alongside another entity or receiving assistance from it. Research on collaborative learning, teamwork, and organisational behaviour has long distinguished collaboration from related but distinct concepts such as cooperation, coordination, and division of labour. What differentiates collaboration is not the presence of multiple actors, but the nature of the relationship among them. Clarifying what constitutes genuine collaboration is therefore a necessary precondition for assessing claims that human-AI relationships can be described in these terms.

Roschelle and Teasley's (1995) foundational work on collaborative problem solving provides a useful starting point. They define collaboration as "a coordinated, synchronous activity that is the result of a continued attempt to construct and maintain a shared conception of a problem" (p. 70). Central to this definition is the notion of a joint problem space: a shared understanding that participants actively build, monitor, and revise through interaction. Collaboration, in this sense, requires more than working on the same task or contributing separate parts to a whole. It depends on reciprocal engagement, negotiation of meaning, and sustained mutual orientation to a shared problem. Dillenbourg (1999) reinforces this view by arguing that collaboration cannot be reduced to a single behavioural feature, but instead involves the interaction of situational conditions, communicative processes, cognitive and social mechanisms, and learning outcomes. He explicitly distinguishes collaboration from mere co-presence or coordination, emphasising that shared meaning-making is a defining requirement.

Research on collaboration in workplace and organisational settings provides further conceptual clarity. Bedwell et al. (2012), synthesising decades of research, propose an integrative, multilevel conceptualisation of collaboration that foregrounds jointly defined objectives, coordinated and interdependent action, and collective ownership of outcomes. Related work emphasises the role of shared mental models: common understandings of goals, roles, and coordination patterns that enable collaborators to anticipate and adapt to one another's actions (Mathieu et al., 2008). From this perspective, collaboration is best understood as a relational and cognitive achievement rather than a functional arrangement for distributing tasks (Salas et al., 2008).

Synthesising these bodies of work, collaboration can be understood as comprising seven interrelated criteria. First, collaboration involves shared goals and purpose, such that participants hold a mutual understanding of objectives and are jointly committed to achieving them. Second, collaboration requires interdependence and coordination, with participants relying on one another for task completion and aligning their actions over time. Third, collaboration depends on reciprocal communication through which meaning is negotiated and shared understanding is constructed. Fourth, collaboration entails distributed accountability, whereby responsibility for outcomes is shared among participants. Fifth, collaborators contribute complementary expertise that is integrated rather than merely aggregated. Sixth, collaboration involves calibrated trust, developed through demonstrated competence and

reliability. Finally, collaboration includes acknowledged contribution, with participants' inputs recognised and appropriately attributed (Bedwell et al., 2012; Costa et al., 2001; Mathieu et al., 2008; Salas et al., 2008; Thompson, 2008; Wagner, 1995). These criteria overlap and are mutually dependent rather than distinct or independently sufficient. Interdependence without shared purpose, communication without accountability, or complementary expertise without trust may support coordination or assistance, but they do not constitute collaboration as it is defined in the literature.

Work extending collaboration theory to human-AI contexts underscores the importance of maintaining these distinctions. Seeber et al. (2020), drawing on the expertise of collaboration scholars, examine what it would mean for AI systems to function as teammates rather than tools within their “machines as teammates” research agenda. Rather than assuming that AI systems already meet collaborative criteria, they identify conceptual and practical challenges associated with attributing collaboration to non-human actors, particularly in relation to coordination, shared understanding, and accountability.

This conceptualisation provides a framework for examining claims that human-AI use in language education constitutes collaboration. It distinguishes surface similarities, such as task interdependence or complementary capability, from the relational features that define collaboration as a social practice. The sections that follow apply this framework to human-AI use in language education, examining how each criterion is instantiated, constrained, or absent. These criteria are revisited in consolidated form later in the paper to clarify where human–AI use aligns with, and departs from, collaboration as a defined relationship.

Exploring Human–AI Collaboration

Research across collaboration theory and human–computer interaction consistently emphasises that collaboration is not defined by joint activity alone, but by shared intention, reciprocal understanding, accountability, and social embeddedness (Bedwell et al., 2012; Seeber et al., 2020). When the defining features of collaboration are examined more closely, important limitations become apparent. These limitations are not best understood as temporary shortcomings that will be resolved through improved models, larger datasets, or more refined interfaces. Rather, they arise from fundamental differences between human collaborators and AI systems in relation to intention, understanding, accountability, and social embeddedness. Clarifying where and why human-AI interaction fails to meet core collaborative criteria is therefore essential, not to reject AI use in language education, but to understand what kinds of relationships AI can and cannot sustain, and what pedagogical risks arise when collaboration language obscures these distinctions.

It would be misleading to suggest that human-AI interaction shares nothing with collaboration. In practice, teachers, learners, and researchers often organise their work around AI in ways that resemble collaborative activity, particularly where tasks are divided, reliance is mutual, or contributions are combined to produce an outcome. Ignoring these points of alignment would oversimplify how generative AI is currently used in language education. At the same time, the presence of such similarities does not in itself establish that collaboration is taking place in the sense articulated in the preceding section. To assess the collaboration framing more carefully, it is therefore necessary to examine where human-AI interaction aligns with recognised features of collaboration, and to consider what these alignments do and do not imply for how AI functions within educational practice. To assess the collaboration framing more carefully, it is therefore necessary to examine where human–AI interaction aligns with recognised features of collaboration, and where it does not. The sections that follow consider each criterion in turn, not to determine whether human–AI interaction can become collaborative, but to clarify what these apparent alignments actually signify in practice.

Task Interdependence

Perhaps the clearest area of alignment with collaboration is task interdependence, albeit in a constrained and asymmetrical form. When language teachers use AI to generate initial drafts of teaching materials, they may genuinely depend on AI output to progress their work within available time and resource constraints. At the same time, the AI cannot generate pedagogically meaningful output without substantial human input in the form of task definition, contextual framing, and evaluative guidance. The teacher provides purpose, constraints, and judgement, while the AI supplies text generation capacity at a scale and speed the teacher could not achieve alone. This creates a form of conditional interdependence: neither party can produce the outcome independently, even though agency and responsibility remain entirely human.

Similarly, learners who use AI for writing support often organise their work around AI contributions, while simultaneously shaping those contributions through prompts, feedback, and selection. A student may rely on AI to generate vocabulary suggestions, provide structural feedback, or offer alternative phrasings, but the usefulness of such output depends on the learner's ability to specify needs, evaluate suggestions, and integrate them into their own work. Researchers, too, may depend on AI for tasks such as literature synthesis, while remaining responsible for defining scope, judging relevance, and validating outputs. In each case, the task outcome emerges through interaction, even though the relationship lacks shared purpose or reciprocal understanding.

Empirical research supports these observations. Systematic reviews of generative AI use in language education show that writing tasks dominate research attention, with 42.4% of studies focusing on writing-related applications (Wang et al., 2025). This pattern reflects a form of functional interdependence in practice: teachers and learners increasingly rely on AI output, while AI output itself remains dependent on human framing and evaluation, resulting in reliance that is mutual in process but asymmetrical in control and accountability.

Complementary Capabilities

Human-AI interaction also exhibits complementary capabilities, though this complementarity differs in important ways from that found in human collaboration. Whereas task interdependence concerns how work is structured and divided, complementary capabilities concern the nature of what each participant brings to the activity. AI systems offer capacities that humans typically cannot match, including rapid text generation, access to large volumes of information, continuous availability, and tolerance for repetition. Humans, by contrast, contribute contextual understanding, pedagogical judgement, creative direction, and the ability to evaluate output in relation to goals that extend beyond the immediate task. Research on human-AI performance consistently suggests that benefits arise not from substitution, but from the careful combination of distinct human and machine strengths under appropriate task conditions (Vaccaro & Waldo, 2024).

In language education, this contrast is particularly evident. Teachers may use AI to generate multiple versions of instructional materials, suggest adaptations for different proficiency levels, or provide feedback on learner writing at a scale that would be impractical for an individual teacher. Learners may engage with AI for low-stakes practice, immediate feedback, or exposure to alternative linguistic forms outside class hours. Researchers may draw on AI to assist with literature synthesis, idea generation, or drafting, thereby extending their cognitive reach when working with large or complex bodies of material. In each case, AI contributes capacity, while humans retain responsibility for direction, interpretation, and evaluation.

However, this form of complementarity differs from the complementary expertise that characterises collaboration. In collaborative relationships, participants bring distinct but commensurable forms of knowledge that are integrated through shared understanding, mutual adjustment, and joint responsibility for outcomes. By contrast, AI contributes functional capacity without understanding the domain in which it operates or the purposes it serves. Empirical work on human–AI delegation highlights this asymmetry, showing that productive outcomes depend on humans maintaining control over goal setting, judgement, and decision-making rather than deferring these functions to AI systems (Fügener et al., 2022). As a result, the human contribution is not merely complementary but constitutive: without human framing, evaluation, and accountability, AI output has no inherent educational or scholarly value.

This distinction matters for how complementarity is interpreted. While combining human and AI strengths can enhance efficiency and scale, such combinations do not in themselves establish the reciprocal knowledge integration that collaboration implies. Complementary capabilities may support performance, but they do not constitute shared expertise or joint understanding, which remain central to established conceptions of collaboration (Hemmer et al., 2025).

Acknowledged Contribution

Norms around acknowledging AI use are emerging across educational and scholarly contexts. Many journals now require authors to disclose the use of generative AI in manuscript preparation, while institutions increasingly specify how learners should report AI assistance in assessed work (Jackson et al., 2023). In classroom practice, teachers develop rules governing transparency about AI involvement, often distinguishing between permitted and prohibited forms of use. These developments reflect growing recognition that AI-mediated activity needs to be visible rather than hidden, particularly in contexts where evaluation, learning outcomes, or scholarly contribution are at stake.

However, acknowledgement in this sense differs fundamentally from how contributions are recognised in collaborative relationships. In human collaboration, acknowledgement is tied to intellectual contribution, shared ownership, and, in many cases, formal credit such as authorship or co-authorship. Such acknowledgement presupposes agency, responsibility, and the capacity to justify one's contribution. By contrast, acknowledging AI use functions primarily as procedural disclosure rather than attribution of credit. It clarifies how work was produced, but it does not imply shared authorship, ownership, or responsibility for the outcome.

This distinction has important implications for claims about collaboration. Scholarship on authorship and contribution emphasises that acknowledgement signals not only participation, but also accountability and intellectual ownership (Birnholtz, 2006; Allen et al., 2014). AI systems cannot assume responsibility for claims made, cannot defend methodological or interpretive decisions, and cannot be held accountable for errors or ethical breaches. As a result, while AI involvement may be acknowledged for reasons of transparency or integrity, it cannot be credited in the way human collaborators are. The act of acknowledgement therefore reinforces, rather than dissolves, the asymmetry between human and AI contributions.

Shared Goals and Purpose

Collaboration requires that partners hold mutual understanding of objectives and joint commitment to achieving them. Research on teamwork and collaborative work stresses that shared goals involve intentional orientation toward common ends, responsiveness to those ends over time, and the capacity to adjust one's contributions in light of collective purpose (Mathieu et al., 2008; Salas et al., 2008). When two teachers collaborate on curriculum development, both understand what they are trying to

accomplish, care about the outcome, and adjust their contributions in service of shared goals. This mutual purpose shapes how each participant engages with the task.

AI systems do not hold purposes. When a language teacher uses ChatGPT to develop a lesson plan, the teacher brings purpose, such as promoting learner engagement, addressing specific learning objectives, or fitting institutional constraints. The AI generates text that is statistically plausible given its training data and the prompt provided, but it does not understand educational goals as goals, does not value outcomes, and cannot orient its behaviour toward shared ends. This distinction has been repeatedly emphasised in HCI and AI ethics research, which cautions against attributing intention, commitment, or goal-directed agency to computational systems on the basis of surface-level behavioural fluency (Suchman, 2007; Floridi et al., 2018; Birhane et al., 2023). This difference places clear limits on what can reasonably be described as shared purpose in human-AI work.

This asymmetry has practical consequences. Human collaborators can be asked why they made particular choices and can justify decisions by reference to shared goals. AI systems can generate explanations, but these are post-hoc rationalisations rather than accounts of genuine reasoning or commitment (Bender et al., 2021). Teachers, learners, and researchers working with AI carry the full burden of purpose; AI contributes capacity without commitment.

Reciprocal Communication and Sensemaking

Genuine collaboration involves not merely exchanging information but jointly constructing understanding. Swain's (2000, 2006) work on collaborative dialogue in second language acquisition demonstrates how learners working together engage in "languageing"—using language to work through problems and construct meaning. This process requires mutual understanding: each participant must comprehend not just the words but the intentions, uncertainties, and reasoning of the other.

Interaction with AI resembles dialogue but lacks the reciprocal understanding that makes dialogue generative. When a learner prompts AI for feedback on writing, the AI processes the text and generates a response based on pattern matching against training data. The learner may refine their prompt based on the response, creating an iterative exchange. However, this is not joint sense-making. The AI does not understand the learner's communicative intentions, does not recognise when its response misses the point, and cannot genuinely negotiate meaning. Similar concerns are well documented in HCI research on conversational agents, which distinguishes conversational fluency from understanding and warns against equating interactional form with communicative substance (Luger & Sellen, 2016; Suchman, 2007). The result is interaction without shared understanding, which places clear limits on how far such exchanges can be considered collaborative.

The limitations become apparent when communication fails. Human collaborators can recognise misunderstanding, ask clarifying questions from genuine uncertainty, and repair breakdowns by reference to shared context. AI can respond to correction, but it does not understand what went wrong, as it simply generates new output conditioned on additional input. For language learners, this distinction is critical, as negotiation of meaning has long been identified as a central mechanism supporting language development (Long, 1996). What appears interactionally similar therefore rests on very different underlying processes, with important implications for how communicative activity supports learning.

Distributed Accountability

Perhaps the clearest failure concerns accountability. Collaboration distributes responsibility among partners, each of whom can be held accountable for their contributions (Bedwell et al., 2012). When

co-authors produce a flawed publication, both bear responsibility. When team members fail to complete assigned tasks, consequences follow. AI cannot be held accountable. When AI-generated teaching materials contain errors, the teacher bears full responsibility. When AI-assisted student writing includes fabricated citations—a well-documented phenomenon with large language models—the student faces academic integrity consequences. When AI-assisted research produces methodological problems, the researcher answers to reviewers, institutions, and the scholarly community.

AI systems cannot be held accountable. When AI-generated teaching materials contain errors, responsibility rests with the teacher. When AI-assisted student writing includes fabricated citations, the student faces academic integrity consequences. When AI-assisted research produces methodological problems, the researcher answers to reviewers, institutions, and the scholarly community. Research on automation and responsibility consistently demonstrates that accountability in human–machine systems remains unavoidably human, regardless of how labour is distributed (Mittelstadt et al., 2016; Johnson & Verdicchio, 2017).

This asymmetry is not a temporary limitation but a structural feature. Accountability requires capacity for consequences: professional sanctions, reputational damage, remedial obligations. AI systems cannot experience consequences in any meaningful sense. The human in a human–AI interaction remains solely accountable regardless of how the work was divided. The accountability gap has particular significance for language education. Stockwell (2024) describes an “AI paradox” where teachers create tasks for students using AI that students complete using AI, obscuring who is responsible for learning outcomes. While academic integrity frameworks place responsibility entirely on learners, collaboration language implies shared ownership that does not align with actual accountability structures.

Calibrated Trust

Lee and See’s (2004) influential framework on trust in automation emphasises the importance of calibrated trust, that is, relying on automated systems appropriately based on their actual capabilities and limitations. Trust that exceeds system reliability leads to misuse; trust that falls short leads to disuse. Appropriate reliance requires understanding what a system can and cannot do. The trust calibration problem with generative AI is particularly acute because these systems perform variably across domains and tasks in ways that are difficult to predict. AI may produce highly competent output for one language education task while generating plausible but incorrect output for another, with no obvious indicators of reliability. This creates an expertise paradox: accurately calibrating trust in AI output requires the very expertise that AI is meant to supplement (Amershi et al., 2019). Empirical research on decision-support systems shows that users under uncertainty are often inclined to defer to AI recommendations even when those recommendations are flawed, particularly when they lack domain expertise (Gaube et al., 2021). As a result, appropriate reliance on AI remains a demanding human judgement rather than a shared or distributed responsibility.

The consequences fall most heavily on those with least expertise. Language learners, still developing target language competence, are poorly positioned to evaluate AI-generated language output and may accept errors as correct, adopt inappropriate register, or fail to recognise when AI suggestions contradict sound learning practices. Teachers and researchers face similar challenges when AI operates outside their areas of expertise. Research on AI overreliance demonstrates that users frequently overestimate AI capabilities and accept outputs uncritically, particularly when they lack domain expertise to evaluate accuracy (Wang et al., 2023). This excessive reliance can degrade critical thinking skills and decision-making abilities, precisely the capacities language education seeks to develop. Furthermore,

trust in collaboration is bidirectional, where partners develop mutual trust through demonstrated reliability. AI systems do not trust their human users as the concept is not comprehended by them. The one-directional nature of trust in human-AI interaction represents another departure from collaborative relationships.

Perspectives from Teaching, Learning, and Research

The limitations identified in the discussion above apply broadly to human-AI interaction in language education, but they are not experienced uniformly. Teachers, learners, and researchers occupy different positions in relation to purpose, responsibility, evaluation, and learning, and these positions shape how AI use is interpreted and enacted in practice. Examining these perspectives separately helps to clarify why collaboration framing is not merely a theoretical concern, but one with concrete pedagogical, ethical, and professional consequences. While all three groups interact with the same underlying technologies, they do so under different expectations, constraints, and forms of accountability. As a result, the risks associated with over-attributing collaboration to human-AI interaction manifest in distinct ways across teaching, learning, and research contexts. Looking at these role-specific experiences therefore sheds light on how abstract limitations translate into everyday practice, and why precision in how human-AI relationships is described is crucial.

Shared Challenges

All three groups—teachers, learners, and researchers—face a common verification burden. Confirming that AI output meets appropriate quality standards requires expertise and time; for teachers reviewing AI-generated materials, learners evaluating AI feedback, and researchers checking AI-assisted writing, this verification effort may approach or even exceed the effort required to produce the output independently. This burden falls entirely on the human, as AI systems cannot verify the accuracy, appropriateness, or consequences of their own output in any meaningful sense. From a critical AI literacy perspective, this highlights that effective AI use depends not on trust in outputs, but on users' capacity to interrogate, contextualise, and evaluate them (Long & Magerko, 2020; Selwyn, 2022).

All three groups also confront challenges related to skill maintenance. When AI handles tasks through which expertise is ordinarily developed, reliance on AI risks undermining the acquisition and retention of capabilities that remain central to teaching, learning, and research. Stockwell (2024) draws an analogy to calculators: while most people can perform complex calculations with technological assistance, many struggle to do so unaided, even for relatively simple arithmetic. Similar dynamics may affect writing, pedagogical design, and research practices when users consistently offload cognitively demanding tasks to AI systems. Research on AI literacy emphasises that understanding when not to rely on AI is as important as knowing how to use it, particularly in educational contexts where learning depends on productive effort and reflective engagement (Ng et al., 2021; Kasneci et al., 2023).

Questions of attribution and acknowledgement arise from these same structural conditions. Because AI output must be verified, interpreted, and integrated by humans, responsibility for the final product remains entirely human, even when AI involvement is substantial. Critical AI literacy frameworks stress the importance of recognising how power, responsibility, and accountability are distributed in human-AI systems, rather than treating AI contributions as neutral or autonomous (Selwyn, 2022; Williamson & Eynon, 2020). Teachers, learners, and researchers must therefore navigate attribution practices that require transparency without implying shared authorship or responsibility, reinforcing the limits of collaboration framing in everyday practice.

Distinctive Teacher Concerns

Language teachers carry professional responsibility for pedagogical decision-making in ways that cannot be redistributed to a generative system. Even when AI produces plausible lesson materials, explanations, or feedback at scale, the teacher remains accountable for alignment with learning objectives, suitability for particular learners, and consistency with curricular and assessment requirements. This reflects a long-standing distinction in educational technology research between tools that extend instructional capacity and the professional judgement required to decide how, when, and why such tools should be used (Selwyn, 2023). In this respect, AI may alter how pedagogical work is carried out, but it does not alter where responsibility ultimately resides.

A further complication is that teachers' use of generative AI is often shaped by contextual pressures that make disclosure and critical scrutiny uneven. Recent research suggests that some teachers adopt AI strategically while simultaneously managing concerns about professional legitimacy, trust, and identity, including decisions about whether and how to reveal AI use in their own teaching practice (Barnes & Tour, 2025; Choi et al., 2025)). This matters for collaboration framing because describing teacher–AI interaction as collaboration risks implying shared ownership of pedagogical work, when responsibility for pedagogical decisions and their consequences remains firmly located with the teacher.

There is also an expertise-related dimension that appears collaborative on the surface but functions differently in practice. Research with language teacher educators indicates that generative AI is likely to reshape what needs to be taught and assessed in initial teacher education, while also exposing uneven confidence and preparedness to address AI-related challenges (Moorhouse & Kohnke, 2024; Kohnke, 2025). As AI produces competent-seeming output across a wide range of domains, teachers are increasingly required to evaluate materials at the edges of their expertise, where error detection is most difficult. This expands the verification burden and reduces opportunities for the kind of productive struggle through which professional knowledge is developed.

Work on AI literacy in teacher education reinforces this asymmetry. AI literacy is not limited to technical competence but includes professional judgement about appropriate use, evaluation of outputs, and modelling responsible practices for learners (Sperling et al., 2024). In this sense, framing teacher–AI interaction as collaboration risks obscuring the very asymmetry that matters most: teachers remain the locus of judgement, responsibility, and accountability, even when AI is heavily involved in production. For teachers, then, describing AI use as collaboration risks obscuring where pedagogical judgement, responsibility, and accountability ultimately reside, with consequences for how teaching decisions are justified and modelled for learners.

Distinctive Learner Concerns

Learners face a particularly complex relationship with AI because the central aim of language education is the development of communicative competence rather than the production of acceptable text. When AI supplies fluent language that learners incorporate into their work, observable performance may improve even if underlying competence remains unchanged. This distinction is well established in second language research, yet generative AI intensifies it by making high-quality output available without requiring equivalent internalisation or decision-making.

This tension aligns with research on productive struggle in learning. Effortful engagement, error-making, and iterative refinement are not obstacles to language acquisition but key mechanisms through which competence develops. When AI provides ready-made solutions by suggesting formulations, correcting errors, or restructuring text, it may reduce opportunities for learners to engage in these

processes. Barrot (2023) characterises generative AI in L2 writing as offering both pedagogical potential and significant risk, noting that while learners often value AI support and produce more polished texts, reliance on AI can lead to shallow processing and weakened engagement with the learning task itself.

Learners also face particular difficulty calibrating trust in AI output because they are still developing the linguistic and pragmatic competence needed to evaluate appropriateness and accuracy. Research examining learner and instructor perspectives on ChatGPT use in EFL contexts highlights widespread uncertainty about permissible use and responsibility, particularly in relation to academic integrity (Almanea, 2024). In such contexts, collaboration framing may further blur perceived boundaries by suggesting shared ownership of output, making it harder for learners to maintain a clear sense of what must remain demonstrably their own work. For learners, framing AI use as collaboration risks confusing improved performance with learning itself, potentially undermining the development of communicative competence by masking where effort, judgement, and responsibility for meaning-making should lie.

Distinctive Researcher Concerns

Researchers' use of AI raises particularly complex questions about intellectual contribution and scholarly responsibility. When AI assists with drafting, synthesising literature, or shaping the expression of ideas, it becomes increasingly difficult to specify what constitutes the researcher's original contribution. Publication ethics frameworks are grounded in assumptions of human authorship, intention, and accountability, and these assumptions do not translate cleanly to AI-assisted research practices.

This difficulty is compounded by challenges surrounding methodological transparency. Researchers are now commonly required to disclose AI use, yet existing conventions often provide little guidance on how detailed such disclosures should be. Minor uses of AI may be too routine to document meaningfully, while more substantial uses raise questions about whether AI has influenced argumentation, interpretation, or claims of originality. Moorhouse et al. (2025) argue that without clearer disciplinary guidance, transparency risks becoming a procedural formality rather than a substantive account of how research was conducted.

Peer review introduces further complications. Reviewers may be asked to evaluate work in which the boundaries of human intellectual labour are unclear, particularly when AI-generated text is fluent and persuasive. This uncertainty does not simply concern detection, but evaluation: peer review presumes that authors can explain, defend, and take responsibility for the reasoning underlying their claims. The problem is intensified by the documented tendency of generative AI systems to produce plausible but fabricated references, increasing the verification burden placed on reviewers and editors.

These issues underscore why collaboration framing is especially problematic in research contexts. Collaboration implies shared intellectual ownership and distributed responsibility, yet AI systems cannot justify methodological decisions, respond to critique, or bear consequences for error. For researchers, precision in how AI involvement is framed therefore functions as a safeguard for scholarly standards, rather than as a terminological preference. For researchers, framing AI as a collaborator risks blurring intellectual ownership and accountability, threatening the assumptions about authorship, originality, and responsibility on which scholarly evaluation depends.

Synthesising the Criteria

The preceding sections have examined how each criterion commonly associated with collaboration is realised, partially realised, or structurally absent in human–AI use in language education. Bringing

these criteria together makes it possible to see patterns that are less visible when they are considered in isolation. Table 1 summarises the concept of human-AI collaboration in language education, highlighting where surface similarities to collaboration emerge and where fundamental limitations remain.

Table 1 *Collaboration Criteria and their Realisation in Human–AI Use in Language Education*

Criterion associated with collaboration	Realisation in human–AI use	Details
Interdependence and coordination	Asymmetrical	Humans may depend on AI output to complete tasks, but coordination is one-sided and human-directed.
Complementary expertise	Partial	AI contributes capacity (e.g. speed, scale), while humans provide judgement and evaluation.
Acknowledged contribution	Procedural	AI use may be disclosed for transparency, but does not entail shared ownership or credit.
Shared goals and purpose	Absent	Goals, values, and commitments originate with the human; AI does not hold or pursue purposes.
Reciprocal communication	Asymmetrical	Interaction is iterative and responsive, but lacks shared understanding or negotiated meaning.
Distributed accountability	Absent	Responsibility for outcomes remains entirely human; AI cannot be held accountable.
Calibrated trust	Asymmetrical	Trust operates in one direction only and requires ongoing human calibration.

The criteria in Table 1 show that human–AI use in language education aligns with collaboration unevenly and asymmetrically. Some features, such as interdependence and complementary capabilities, are readily observable in practice, while others, including shared purpose, reciprocal understanding, and distributed accountability, are systematically absent. This uneven alignment is important because collaboration is not defined by the presence of any single feature, but by how these features operate together. Where key elements are missing, the relationship shifts from collaboration toward other forms of support or assistance, even when interaction appears fluent or productive. Making these distinctions explicit helps clarify why collaboration language can be tempting, yet ultimately misleading, when applied to human–AI relationships.

Reconsidering Human–AI Collaboration

This paper suggests that human-AI use in language education does not meet several criteria that are central to genuine collaboration. Certain features of collaboration appear to be present, most notably task interdependence, complementary capabilities, and emerging norms of acknowledgement. However, these similarities remain partial and surface-level. AI systems do not share purposes, cannot engage in reciprocal understanding, bear no accountability for outcomes, and do not participate in bidirectional trust. As a result, describing human-AI use as collaboration risks overstating what is actually shared between human and machine.

This is not just a terminological distinction. Collaboration is a strong concept, carrying assumptions about shared agency, joint commitment, and distributed responsibility. When AI is framed as a collaborator, these assumptions are implicitly extended to systems that cannot sustain them. Such framing may reduce vigilance in verification, suggest shared responsibility where none exists, or encourage evaluation of learning in terms that prioritise output quality over competence development. Returning

explicitly to collaboration here helps clarify that the issue is not whether AI is useful, but whether collaboration is the right lens through which to understand its role.

This concern resonates with earlier debates in technology-enhanced language learning. Bax's (2003, 2011) work on the normalisation of CALL argued for technologies to become integrated into practice without exaggerated reverence or fear, while maintaining clarity about their pedagogical function. Importantly, normalisation did not imply equivalence between human and technological contributions. The emergence of generative AI raises a parallel challenge: integration must not come at the cost of conceptual inflation, where the language of collaboration obscures meaningful differences in agency and responsibility.

From this perspective, alternative ways of describing human-AI use do not replace collaboration so much as delineate its limits. Framing AI as a **tool** brings human agency and control to the foreground, moving purpose and accountability to the user while recognising that tools can extend capacity. Describing AI as a **resource** highlights the need for selection, evaluation, and critical judgement, reminding users that outputs require scrutiny rather than blind acceptance. Referring to AI as an **interlocutor** acknowledges the interactive and dialogic form of engagement without implying without attributing agency, intention, or shared understanding. Each of these framings preserves the analytic boundary that collaboration requires.

Maintaining this boundary has practical implications across language education. For teachers, it reinforces that pedagogical judgement and responsibility remain human, even when AI contributes substantially to lesson design or feedback. For learners, it helps distinguish assistance from learning itself, reducing the risk of equating improved output with developing competence. For researchers, it clarifies expectations around intellectual ownership, accountability, and methodological transparency. In each case, treating collaboration as a clearly defined concept limits how far it can be extended to human-AI relationships.

Future research can build on this clarification by examining how different framings of human-AI use shape actual practices and outcomes in language education. Empirical work might investigate whether describing AI as a tool, resource, or interlocutor influences learner engagement, reliance patterns, or perceptions of responsibility, as well as how teachers' framing choices affect task design and assessment. Longitudinal studies following learners' and teachers' AI use over time would be particularly valuable for understanding how these framings interact with the development of linguistic competence, pedagogical judgement, and AI literacy.

References

- Allen, L., Scott, J., Brand, A., Hlava, M., & Altman, M. (2014). Publishing: Credit where credit is due. *Nature*, 508(7496), 312–313. <https://doi.org/10.1038/508312a>
- Almanea, M. (2024). Instructors' and learners' perspectives on using ChatGPT in English as a foreign language courses and its effect on academic integrity. *Computer Assisted Language Learning*. <https://doi.org/10.1080/09588221.2024.2410158>
- Amershi, S., Weld, D., Vorvoreanu, M., et al. (2019). Guidelines for human-AI interaction. Proceedings of the CHI Conference on Human Factors in Computing Systems, 1–13. <https://doi.org/10.1145/3290605.3300233>
- Barnes, M., & Tour, E. (2025). Teachers' use of generative AI: A "dirty little secret"? *Language and Education*. <https://doi.org/10.1080/09500782.2025.2485935>
- Barrot, J. S. (2023). Using ChatGPT for second language writing: Pitfalls and potentials. *Assessing Writing*, 57, 100745. <https://doi.org/10.1016/j.asw.2023.100745>

- Bax, S. (2003). CALL—past, present and future. *System*, 31(1), 13–28. [https://doi.org/10.1016/S0346-251X\(02\)00071-4](https://doi.org/10.1016/S0346-251X(02)00071-4)
- Bax, S. (2011). Normalisation revisited: The effective use of technology in language education. *International Journal of Computer-Assisted Language Learning and Teaching*, 1(2), 1–15. <https://doi.org/10.4018/ijcallt.2011040101>
- Bedwell, W. L., Wildman, J. L., DiazGranados, D., Salazar, M., Kramer, W. S., & Salas, E. (2012). Collaboration at work: An integrative multilevel conceptualization. *Human Resource Management Review*, 22(2), 128–145. <https://doi.org/10.1016/j.hrmr.2011.11.007>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Birnholtz, J. P. (2006). What does it mean to be an author? The intersection of credit, contribution, and collaboration in science. *Journal of the American Society for Information Science and Technology*, 57(13), 1758–1770. <https://doi.org/10.1002/asi.20380>
- Choi, K. Y., Wu, C., & Moorhouse, B. L. (2025). Exploring the use of generative artificial intelligence (GenAI) in English language teaching: Voices from in-service teachers at an early-adopting Hong Kong secondary school. *Technology in Language Teaching & Learning*, 7(3), 102516. <https://doi.org/10.29140/tlsl.v7n3.102516>
- Costa, A. C., Roe, R. A., & Taillieu, T. (2001). Trust within teams: The relation with performance effectiveness. *European Journal of Work and Organizational Psychology*, 10(3), 225–244. <https://doi.org/10.1080/13594320143000654>
- Dillenbourg, P. (1999). *Collaborative learning: Cognitive and computational approaches*. Pergamon.
- Floridi, L., Cows, J., Beltrametti, M., et al. (2018). AI4People—An ethical framework for a good AI society. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Fügener, A., Grahl, J., Gupta, A., & Ketter, W. (2022). Cognitive challenges in human–artificial intelligence collaboration: Investigating the path toward productive delegation. *Information Systems Research*, 33(2), 678–696. <https://doi.org/10.1287/isre.2021.1079>
- Gaube, S., Suresh, H., Raue, M., Merritt, A., Berkowitz, S. J., Lerner, E., Coughlin, J. F., Guttag, J. V., Colak, E., & Ghassemi, M. (2021). Do as AI say: Susceptibility in deployment of clinical decision-aids. *npj Digital Medicine*, 4, Article 31. <https://doi.org/10.1038/s41746-021-00385-9>
- Hemmer, P., Schemmer, M., Kühl, N., Vössing, M., & Satzger, G. (2025). Complementarity in human-AI collaboration: Concept, sources, and evidence. *European Journal of Information Systems*, 34(6), 979–1002. <https://doi.org/10.1080/0960085X.2025.2475962>
- Hollan, J., Hutchins, E., & Kirsh, D. (2000). Distributed cognition: Toward a new foundation for human–computer interaction research. *ACM Transactions on Computer-Human Interaction*, 7(2), 174–196. <https://doi.org/10.1145/353485.353487>
- Huxham, C., & Vangen, S. (2005). *Managing to collaborate: The theory and practice of collaborative advantage*. Routledge.
- Jackson, J., Landis, G., Baskin, P. K., Hadsell, K. A., & English, M. (2023). CSE guidance on machine learning and artificial intelligence tools. *Science Editor*, 46(2). <https://doi.org/10.36591/SE-D-4602-07>
- Johnson, D. G., & Verdicchio, M. (2017). Reframing AI discourse. *Minds and Machines*, 27(4), 575–590. <https://doi.org/10.1007/s11023-017-9417-6>
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Krusche, S., Kühl, N., Michels, J., 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. *Educational Psychology Review*, 35, Article 102. <https://doi.org/10.1007/s10648-023-09733-5>
- Kittur, A., et al. (2013). The future of crowd work. *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems*, 1301–1318. <https://doi.org/10.1145/2441776.2441923>

- Kohnke, L. (2025). Empowering pre-service teachers with generative artificial intelligence and micro-learning: Pathways to self-directed growth. *Australian Journal of Applied Linguistics*, 8(4), 102889. <https://doi.org/10.29140/ajal.v8n4.102889>
- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- Long, M. H. (1996). The role of the linguistic environment in second language acquisition. In W. C. Ritchie & T. K. Bhatia (Eds.), *Handbook of second language acquisition* (pp. 413–468). Academic Press.
- Luger, E., & Sellen, A. (2016). “Like having a really bad PA”: The gulf between user expectation and experience of conversational agents. *Proceedings of the CHI Conference on Human Factors in Computing Systems*, 5286–5297. <https://doi.org/10.1145/2858036.2858288>
- Lund, B. D., & Wang, T. (2023). ChatGPT and the future of academic integrity. *British Journal of Educational Technology*, 54(5), 1741–1750. <https://doi.org/10.1111/bjet.13306>
- Malhotra, A., & Majchrzak, A. (2014). Managing crowds in innovation challenges. *California Management Review*, 56(4), 103–123. <https://doi.org/10.1525/cmr.2014.56.4.103>
- Mathieu, J., Maynard, M. T., Rapp, T., & Gilson, L. (2008). Team effectiveness 1997–2007: A review of recent advancements and a glimpse into the future. *Journal of Management*, 34(3), 410–476. <https://doi.org/10.1177/0149206308316061>
- Moorhouse, B. L., & Kohnke, L. (2024). The effects of generative AI on initial language teacher education: The perceptions of teacher educators. *System*, 122, 103290. <https://doi.org/10.1016/j.system.2024.103290>
- Moorhouse, B. L., Nejadghanbar, H., & Yeo, M. A. (2025). Study quality in the age of AI: A disciplinary framework for using GenAI in TESOL research. *TESOL Quarterly*. Advance online publication. <https://doi.org/10.1002/tesq.70026>
- Roschelle, J., & Teasley, S. D. (1995). The construction of shared knowledge in collaborative problem solving. In C. O'Malley (Ed.), *Computer supported collaborative learning* (pp. 69–97). Springer. https://doi.org/10.1007/978-3-642-85098-1_5
- Salas, E., Cooke, N. J., & Rosen, M. A. (2008). On teams, teamwork, and team performance: Discoveries and developments. *Human Factors*, 50(3), 540–547. <https://doi.org/10.1518/001872008X288457>
- Salas, E., Sims, D. E., & Burke, C. S. (2008). Is there a “Big Five” in teamwork? *Small Group Research*, 36(5), 555–599. <https://doi.org/10.1177/1046496405277134>
- Seeber, I., Bittner, E., Briggs, R. O., de Vreede, T., de Vreede, G.-J., Elkins, A., Maier, R., Merz, A. B., Oeste-Reiß, S., Randrup, N., Schwabe, G., & Söllner, M. (2020). Machines as teammates: A research agenda on AI in team collaboration. *Information & Management*, 57(2), 103174. <https://doi.org/10.1016/j.im.2019.103174>
- Selwyn, N. (2023). *Education and technology: Key issues and debates* (3rd ed.). Bloomsbury Academic.
- 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>
- Silverman, B. G. (1992). Human–computer collaboration. *Human–Computer Interaction*, 7(2), 165–196. https://doi.org/10.1207/s15327051hci0702_2
- Sperling, K., Stenberg, C.-J., McGrath, C., Åkerfeldt, A., Heintz, F., & Stenliden, L. (2024). In search of artificial intelligence (AI) literacy in teacher education: A scoping review. *Computers and Education Open*, 6, 100169. <https://doi.org/10.1016/j.caeo.2024.100169>
- Stockwell, G. (2024). ChatGPT in language teaching and learning: Exploring the road we’re traveling. *Technology in Language Teaching & Learning*, 6(1), 2273. <https://doi.org/10.29140/tltl.v6n1.2273>
- Suchman, L. A. (1987). *Plans and situated actions: The problem of human–machine communication*. Cambridge University Press.

- Suchman, L. A. (2007). *Human–machine reconfigurations: Plans and situated actions* (2nd ed.). Cambridge University Press.
- Swain, M. (2000). The output hypothesis and beyond: Mediating acquisition through collaborative dialogue. In J. P. Lantolf (Ed.), *Sociocultural theory and second language learning* (pp. 97–114). Oxford University Press.
- Swain, M. (2006). Linguaging, agency and collaboration in advanced second language proficiency. In H. Byrnes (Ed.), *Advanced language learning: The contribution of Halliday and Vygotsky* (pp. 95–108). Continuum.
- Thompson, L. L. (2008). *Making the team: A guide for managers* (3rd ed.). Pearson Prentice Hall.
- Tlili, A., Shehata, B., Adarkwah, M. A., Bozkurt, A., Hickey, D. T., Huang, R., Agyemang, B., & Burgos, D. (2023). What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Computers & Education*, 196, 104699. <https://doi.org/10.1016/j.compedu.2023.104699>
- Vaccaro, M., & Waldo, J. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behaviour*. <https://doi.org/10.1038/s41562-024-02024-1>
- Wagner, J. A. (1995). Studies of individualism–collectivism: Effects on cooperation in groups. *Academy of Management Journal*, 38(1), 152–173. Retrieved from <https://www.jstor.org/stable/256731>
- Wang, T., Lund, B. D., Marengo, A., Pagano, A., Mannuru, N. R., Teel, Z. A., & Pange, J. (2023). Exploring the potential impact of artificial intelligence (AI) on international students in higher education: Generative AI, chatbots, analytics, and international student success. *Applied Sciences*, 13(11), 6716. <https://doi.org/10.3390/app13116716>
- Wang, Y. (2024). Cognitive and sociocultural dynamics of self-regulated use of machine translation and generative AI tools in academic EFL writing. *System*, 126, 103505. <https://doi.org/10.1016/j.system.2024.103505>
- Wang, Y., Zhang, T., Yao, L., & Seedhouse, P. (2025). A scoping review of empirical studies on generative artificial intelligence in language education. *Innovation in Language Learning and Teaching*. Advance online publication. <https://doi.org/10.1080/17501229.2025.2509759>