



Research Papers in Language Teaching and Learning

Vol. 16, No. 1, March 2026, 96-105

ISSN: 1792-1244

Available online at <http://rpltl.eap.gr>

This article is issued under the [Creative Commons License Deed. Attribution](#)

[3.0 Unported \(CC BY 3.0\)](#)

Generative Artificial Intelligence in Language Education: Ethical Dimensions and the Equitable AI in Language Education Model

Eirini Ioanna Delmadorou

Generative Artificial Intelligence (GenAI) is reshaping language education by expanding opportunities for personalized feedback, materials development, and learner support. At the same time, its pedagogical promise is accompanied by significant ethical concerns involving bias, equity, transparency, and inclusivity. This paper critically examines these concerns through the lenses of sociocultural theory, second language acquisition (SLA), and constructivism, treating GenAI as a powerful but non-neutral cultural tool. The review highlights both the benefits of GenAI, such as increased access to adaptive input and output opportunities, and its risks, including algorithmic bias, Anglocentric dominance, and digital inequality. Building on this analysis, the paper introduces the *Equitable AI in Language Education Model* (EALEM), a conceptual framework organized around four guiding principles: inclusivity, transparency, human-in-the-loop mediation, and participatory design. EALEM offers a practical map for integrating GenAI in ways that support linguistic diversity, foster intercultural competence, and promote educational justice. By articulating this framework, the paper contributes to current debates on AI in education and proposes theoretically grounded, pedagogically relevant guidance for the responsible and equitable use of GenAI in language-learning contexts.

Keywords: GenAI, Language education ethics, Algorithmic bias, Digital equity, Sociocultural mediation, AI literacy

1. Introduction

The emergence of Generative Artificial Intelligence (GenAI) represents one of the most consequential technological developments in contemporary education. In recent years, tools such as large language models have begun to transform how learners engage with texts, how educators design instructional materials, and how institutions conceptualize teaching and learning (Li, 2025; Moorhouse & Wong, 2025). Within language education, GenAI is often presented as a catalyst for pedagogical innovation because it

can generate practice materials, provide rapid feedback, and create expanded opportunities for communicative rehearsal (Lee et.al., 2026; Weng & Fu, 2025). Yet these pedagogical possibilities are inseparable from ethical questions concerning bias, equity, transparency, and cultural representation.

The ethical challenges associated with GenAI extend far beyond familiar concerns about plagiarism or academic misconduct. Scholars have drawn attention to algorithmic bias, digital inequality, and the reinforcement of Anglocentric norms in AI-generated materials (García- López et al., 2025; Gabriel, 2024). These concerns are particularly significant in language education, where linguistic identity, cultural meaning, and communicative legitimacy are central pedagogical issues rather than peripheral ones (Canagarajah, 2012; Pennycook, 2017). When generative systems rely disproportionately on English-dominant training data, they risk reproducing linguistic hierarchies and marginalizing less represented languages, thereby intensifying existing global inequalities (Nyaaba et al., 2024). At the same time, unequal access to advanced AI tools can deepen the digital divide between well-resourced and under-resourced educational settings (OECD, 2023; UNESCO, 2023).

The urgency of these issues is practical as well as theoretical. Educators are increasingly encountering GenAI in their classrooms and must decide how to integrate, regulate, or resist its use (Barnett, 2025; Stokes, 2025). The central challenge is how to balance the efficiency and creativity promised by these tools with the responsibility to cultivate critical thinking, inclusivity, and learner autonomy. This tension raises a series of questions: How can AI-generated materials represent diverse perspectives responsibly? Who is accountable when biased outputs shape learner experiences? How should teacher education evolve in order to prepare practitioners for the ethical dilemmas of AI-mediated classrooms?

Despite the growing volume of scholarship on AI in education, much of the literature remains focused on functionality rather than ethics. Systematic reviews have documented a rapid increase in empirical studies examining the use of GenAI tools, but comparatively fewer studies address fairness, transparency, or equity in educational practice (Chaudhry et al., 2022; Yan et al, 2025). Likewise, policy documents provide important high-level principles, yet they often stop short of offering pedagogically specific guidance for classroom implementation (U.S. Department of Education, 2023; UNESCO, 2024). There is, therefore, a clear need for theoretical contributions that identify ethical issues specific to language education and translate them into usable pedagogical guidance.

Although international frameworks increasingly articulate principles such as fairness, accountability, and transparency, these frameworks generally operate at the policy or systems level. Language education, however, faces distinctive ethical challenges because of its close relationship to culture, identity, power, and linguistic representation. This paper argues that existing AI ethics frameworks do not sufficiently address these language-specific dimensions and therefore do not fully support educators in operationalizing ethical AI use in language classrooms.

To address this gap, the present study introduces the *Equitable AI in Language Education Model* (ELEM), a pedagogically grounded framework that integrates sociocultural theory, SLA, and constructivist perspectives in order to provide classroom-oriented guidance for the ethical use of GenAI in language education. The paper critically examines the ethical dimensions of GenAI, identifies both risks and opportunities, and situates them within broader debates on linguistic justice and digital colonialism. It then proposes ELEM as a conceptual framework for responsible practice built around inclusivity, transparency, human-in-the-loop mediation, and participatory design. In doing so, the article aims to contribute not only to theoretical debate but also to practical, ethically grounded innovation in language education.

2. Theoretical Framework

Sociocultural theory offers a productive lens for understanding the integration of GenAI into language education. Rooted in the work of Vygotsky (1978), sociocultural perspectives conceptualize learning as a socially mediated process in which knowledge is constructed through interaction, guided participation, and the use of cultural tools. From this perspective, GenAI can be understood as a mediational artifact that shapes how learners access, produce, and negotiate language. AI-powered prompts, corrective feedback, and dialogic simulations may extend the learner's zone of proximal development by scaffolding performance in ways broadly consistent with Vygotskian principles (Lantolf & Thorne, 2006). At the same time, sociocultural theory cautions against treating tools as neutral: the assumptions, values, and biases embedded in GenAI systems inevitably influence the learning experiences they mediate.

Complementing this perspective are theories of second language acquisition, particularly Krashen's Input Hypothesis and Swain's Output Hypothesis. Krashen (1982) emphasized the importance of comprehensible input as a driver of acquisition, whereas Swain (1995) highlighted the value of pushed output in consolidating linguistic knowledge. GenAI has the potential to influence both processes. On the one hand, large language models can provide learners with abundant, adaptable input through generated dialogues, explanations, and simplified texts. On the other hand, interactive AI systems can create low-stakes opportunities for language production by prompting learners to write, speak, revise, and reformulate their responses (Moorhouse & Wong, 2025). Yet, as Ellis (2003) suggests, interaction in SLA is not merely the exchange of forms; it is also shaped by negotiation, reciprocity, and meaning-making. For that reason, GenAI should be designed to complement rather than replace interaction with teachers and peers.

Constructivist theories further illuminate the pedagogical significance of GenAI in language education. Constructivism emphasizes the active role of learners in building knowledge through exploration, interpretation, and collaboration. From this perspective, AI tools may support learning when they enable learners to test ideas, receive feedback, and revise their understanding in iterative ways. However, constructivism also reminds us that meaningful learning depends on reflection and agency rather than on passive reception of ready-made answers. The educational value of GenAI, therefore, depends not simply on access to technological outputs but on the quality of the pedagogical tasks in which those outputs are embedded.

To connect these theoretical perspectives to the ethical focus of the paper, four core concepts guide the present analysis: equity, fairness, inclusivity, and transparency. Equity refers to the just distribution of opportunities and support in ways that recognize structural differences in learners' needs and circumstances. Fairness concerns the avoidance of discriminatory outcomes and the responsible treatment of learners across linguistic, cultural, and social groups. Inclusivity involves designing learning environments and resources that reflect diversity and allow meaningful participation. Transparency concerns the intelligibility of AI systems, including the extent to which their purposes, limitations, and decision-making processes can be made understandable to educators and learners. These concepts are interdependent: transparency supports fairness by making bias more visible, and inclusivity strengthens equity by ensuring that historically marginalized learners are not excluded from the benefits of innovation.

By integrating sociocultural theory, SLA, constructivism, and these ethical concepts, this paper situates GenAI within a coherent theoretical framework. Sociocultural theory highlights its status as a cultural tool; SLA clarifies its implications for input, output, and interaction; constructivism underscores learner agency

and reflective engagement; and the ethical concepts of equity, fairness, inclusivity, and transparency provide the normative criteria through which AI integration in language education should be evaluated.

3. Literature Review

3.1 Applications of Generative Artificial Intelligence in Language Education

The integration of GenAI into language education has expanded rapidly, with scholars documenting a wide range of applications that affect both instructional practice and learner experience. One of the most prominent uses of GenAI is the creation of instructional materials. Educators can employ AI tools to generate reading passages, vocabulary activities, grammar exercises, discussion prompts, and differentiated materials tailored to diverse proficiency levels (Lee et al., 2026; Moorhouse & Wong, 2025). This flexibility is particularly attractive in language education, where teachers often need to adapt content to different learner needs, interests, and contexts.

GenAI also supports feedback and formative assessment. Compared with earlier forms of computer-assisted language learning, generative models can provide more nuanced explanations, examples, and reformulations that approximate dialogic feedback. Such tools may help learners notice errors, experiment with language choices, and revise their texts in real time. This can increase opportunities for practice, especially in settings where teachers face large classes or limited time (Moorhouse & Wong, 2025; Weng & Fu, 2025).

A growing body of scholarship situates these applications within broader processes of educational digitalization. Li (2025), in a systematic review of empirical research on GenAI in education, notes that many studies focus on efficiency, personalization, and learner engagement. In language education specifically, GenAI is often praised for lowering barriers to practice by enabling conversational simulation, vocabulary support, and immediate feedback. These developments suggest that GenAI may contribute to more adaptive and responsive language learning environments. However, the same literature also makes clear that pedagogical usefulness alone cannot serve as the sole criterion for adoption; ethical considerations remain central.

3.2 Ethical Concerns

Although the pedagogical potential of GenAI is substantial, its widespread adoption raises a series of ethical concerns that require critical scrutiny. Among the most pressing is algorithmic bias. García-López et al. (2025) argue that GenAI systems can reproduce social and cultural inequalities because they are trained on data that reflect existing imbalances in representation and power. In language education, such bias may appear in stereotypical cultural portrayals, the privileging of dominant varieties of English, or the underrepresentation of minority languages and communicative norms. Gabriel (2024) similarly warns that educational uses of GenAI may widen inequity unless issues of access, representation, and cultural responsiveness are addressed directly.

Transparency and accountability constitute a second major concern. Chaudhry et al. (2022) propose a transparency framework for AI in education and emphasize the importance of making system logic, limitations, and outputs understandable to educators and learners. Yet many widely used GenAI tools operate as opaque “black boxes,” making it difficult to determine how outputs are generated, which sources of bias are embedded in the system, or how errors should be interpreted. Dwivedi et al. (2023)

likewise stress that institutions often prioritize innovation and efficiency more readily than sustained ethical reflection, thereby creating governance gaps in educational AI adoption.

Equity and access form a third area of concern. OECD (2023) and UNESCO (2024) both underline the risk that AI adoption may reinforce existing educational inequalities if access to quality tools, infrastructure, and teacher training remains uneven. In language education, this risk is especially acute because AI tools can become gatekeepers to high-quality input, feedback, and learning support. Well-resourced institutions may be able to integrate these tools effectively, while under-resourced contexts remain excluded from their potential benefits.

Finally, ethical concerns extend to learner autonomy. While GenAI can support learners by offering immediate assistance, explanations, and practice opportunities, it may also foster overreliance if used uncritically. Giannakos (2024) notes that the promise of efficiency can obscure deeper questions about what learners are actually doing cognitively when they rely on AI-generated responses. If learners begin to outsource idea generation, revision, and reflection too readily, GenAI may weaken rather than strengthen the independent and critical capacities that language education seeks to cultivate.

3.3 Policy and Guidelines

The growing awareness of these ethical challenges has led to the development of policy frameworks and institutional guidance. UNESCO (2024) has issued recommendations for the use of generative AI in education and research, highlighting the importance of transparency, accessibility, data governance, and human oversight. Similarly, the U.S. Department of Education (2023) frames AI as both an opportunity and a responsibility, arguing that innovation must be accompanied by protections for equity, safety, and educational quality.

National and institutional responses are also evolving. Cassidy (2024) reports that Australian schools have moved toward more formal guidance on classroom AI use, reflecting an attempt to balance opportunity with risk. Barnett (2025) and Stokes (2025) likewise suggest that educators are increasingly experimenting with AI tools while simultaneously recognizing concerns about cheating, data privacy, bias, and professional preparedness. These accounts reveal that policy development is no longer abstract; it is becoming a practical issue of classroom governance and professional judgment.

At the same time, the literature suggests that policy cannot be limited to top-down regulation. García-López et al. (2025) argue that ethical AI in education requires participatory approaches in which educators and learners help shape how systems are implemented. This is especially important in language education, where local linguistic ecologies, cultural identities, and curricular goals vary substantially across contexts. Overall, the literature depicts a field marked by both enthusiasm and caution: GenAI offers real pedagogical affordances, but its ethical implications remain insufficiently resolved.

4. Discussion

The literature reviewed above indicates that the integration of GenAI into language education cannot be reduced to a simple debate between innovation and resistance. Rather, it reveals a more complex terrain in which pedagogical opportunity and ethical risk are deeply intertwined. Language education is especially sensitive to these tensions because language learning is never merely technical; it involves identity,

culture, legitimacy, and participation in wider social worlds. As a result, the ethical evaluation of GenAI in this field must attend to more than functional effectiveness.

A first point of tension concerns representational bias. Large language models inherit the limitations of their training data and may reproduce dominant cultural assumptions in subtle but consequential ways (García-López et al., 2025; Yan et al., 2025). In language classrooms, such bias may shape which accents appear legitimate, which cultural narratives are treated as normative, and which communicative practices are rendered visible or invisible. From the perspective of translingual practice, this is particularly problematic because language education should expand rather than constrain learners' engagement with linguistic diversity (Canagarajah, 2012). A second issue concerns access and structural inequality. OECD (2023) and UNESCO (2024) show that digital inequality remains a major barrier to equitable AI integration. If high-quality AI tools become concentrated in privileged institutions, then GenAI may intensify rather than alleviate educational disparities. The problem is not only material access to devices and subscriptions, but also access to training, institutional support, and pedagogical guidance. In this sense, AI literacy is itself an equity issue. Third, the educator's role must be reconsidered. As GenAI systems assume some functions traditionally associated with teachers, such as generating materials, responding to learner questions, or suggesting revisions, the educator's work shifts from information delivery toward ethical mediation and pedagogical judgment. This does not diminish the teacher's importance; on the contrary, it increases the need for professional discernment. Educators must evaluate outputs, contextualize them, identify bias, and decide when AI use supports learning and when it undermines it (Richards & Rodgers, 2014; U.S. Department of Education, 2023).

Learner autonomy also requires careful reinterpretation. GenAI can create opportunities for independent practice and immediate support, which may increase confidence and engagement. Yet autonomy in language education does not mean simply working alone with a tool; it means developing the capacity to make informed, reflective, and critical choices. If learners treat AI output as authoritative or become dependent on automated assistance, the apparent independence offered by AI may conceal a loss of intellectual agency. Constructivist perspectives, therefore, suggest that GenAI should be used to stimulate inquiry and revision, not to replace them.

Finally, the literature raises broader concerns about digital colonialism. Nyaaba et al. (2024) argue that AI systems developed primarily in the Global North can export dominant epistemologies, languages, and values to other educational settings. In language education, this risk is especially serious because linguistic diversity is inseparable from cultural recognition and educational justice. The ethical challenge is not simply to add more languages to AI systems, but to ensure that local knowledge, communicative practices, and learner identities are treated as pedagogically meaningful rather than as peripheral deviations from a dominant norm. Taken together, these concerns do not justify rejecting GenAI outright. Rather, they point to the need for an explicit ethical framework capable of translating broad principles into pedagogically meaningful guidance. It is in response to that need that the present paper proposes the *Equitable AI in Language Education Model* (EALEM).

5. Proposed Conceptual Framework

Unlike *Equitable AI in Language Education Model* (EALEM), which is proposed as a language-specific, pedagogically grounded framework for the ethical integration of GenAI. Unlike general AI ethics frameworks that remain at the level of broad normative principles, EALEM is designed to address the particular realities of language education, where issues of identity, representation, interaction, and access are central. The model is grounded in the theoretical perspectives outlined earlier and organized around

four interrelated principles: inclusivity, transparency, human-in-the-loop mediation, and participatory design.

The first principle, *inclusivity*, refers to the intentional design and pedagogical use of GenAI tools in ways that reflect linguistic diversity, cultural plurality, and learner identity. An inclusive AI-mediated language classroom does not treat learners as a homogeneous group or present dominant language varieties as universally normative. Instead, it critically examines whose voices, examples, and communicative practices are represented in AI-generated content. In practical terms, this principle calls for educators to review outputs for cultural stereotyping, expand the range of linguistic examples used in instruction, and adapt AI tasks so that diverse learners can participate meaningfully.

The second principle, *transparency*, addresses the opacity of generative systems. Because many GenAI tools function as black boxes, educators and learners may not understand how outputs are produced or why particular responses appear authoritative (Chaudhry et al., 2022). Within EALEM, transparency involves making the limitations, risks, and provisional nature of AI output explicit. This means discussing hallucinations, bias, dataset limitations, and uncertainty with learners rather than presenting AI-generated text as neutral knowledge. Transparency also requires institutions to adopt tools and policies that permit meaningful scrutiny wherever possible.

The third principle, *human-in-the-loop mediation*, emphasizes that GenAI should support rather than displace human pedagogical judgment. Grounded in sociocultural and constructivist understandings of learning, this principle positions educators as mediators who interpret, contextualize, and evaluate AI use in relation to pedagogical goals. Human oversight is especially important in language education because feedback, correction, and cultural framing are not value-free processes. Under EALEM, teachers remain responsible for deciding when AI use is educationally appropriate, how outputs should be discussed, and how learner reflection can be sustained.

The fourth principle, *participatory design*, extends ethical reflection beyond classroom practice to the broader processes through which AI tools are selected, adapted, and governed. García-López et al. (2025) argue that ethical AI in education requires the involvement of those who are affected by its implementation. In the context of language education, this means that educators, learners, and communities should have a voice in how AI systems are integrated, evaluated, and refined. Participatory design is particularly important in multilingual and culturally diverse settings, where top-down technological solutions may overlook local educational priorities and linguistic realities.

Together, these four principles form a framework that is both theoretically grounded and practically adaptable. Inclusivity ensures representation and cultural responsiveness; transparency supports accountability and critical understanding; human-in-the-loop mediation preserves pedagogical judgment; and participatory design promotes legitimacy, responsiveness, and context sensitivity. EALEM does not offer a universal formula, but it provides a structured way of thinking about ethical integration that can guide educators, institutions, and researchers. The implications of EALEM extend beyond individual classrooms. In teacher education, the framework suggests the need to prepare educators not only in technical AI literacy but also in ethical reasoning, bias recognition, and reflective task design. In institutional policy, it points toward more context-sensitive forms of governance that connect general principles to local pedagogical realities. In research, it offers a conceptual basis for examining how ethical AI integration can be operationalized and evaluated across different language-learning contexts.

EALEM also has limitations. Implementing inclusivity may be constrained by the availability of culturally and linguistically diverse datasets. Transparency is difficult when commercial systems do not disclose their underlying processes. Human-in-the-loop mediation requires time, expertise, and institutional support that may not always be available. Participatory design, while desirable, depends on genuine collaboration between educators, learners, developers, and policymakers. For these reasons, EALEM should be understood not as a final solution but as a starting point for continued dialogue, critical reflection, and empirical refinement.

6. Conclusion

The integration of GenAI into language education is both an opportunity and an ethical test. This paper has argued that although GenAI can support personalization, feedback, and expanded access to language practice, its adoption also raises serious questions about bias, equity, opacity, and learner autonomy (Lee et al., 2026; Moorhouse & Wong, 2025). Through the lenses of sociocultural theory, SLA, and constructivism, the discussion has shown that GenAI functions as a mediational tool that reshapes input, output, interaction, and the conditions under which language learning takes place (Krashen, 1982; Swain, 1995; Vygotsky, 1978).

The proposed *Equitable AI in Language Education Model* (EALEM) responds to these challenges by articulating four guiding principles: inclusivity, transparency, human-in-the-loop mediation, and participatory design. Taken together, these principles provide a conceptual map for the ethically grounded adoption of GenAI in language education. They also reframe the role of educators as ethical mediators who must help learners engage critically with AI-generated content rather than consume it unreflectively (Richards & Rodgers, 2014; UNESCO, 2024).

The broader implications of this argument are pedagogical, institutional, and political. Pedagogically, language classrooms must cultivate critical AI literacy alongside linguistic competence. Institutionally, schools and universities must invest in professional learning, infrastructure, and policy guidance that support equitable implementation. Politically, policymakers must recognize that AI in education is not only a matter of innovation, but also of justice. Without deliberate safeguards, GenAI may reproduce digital colonialism and deepen global inequities (Nyaaba et al., 2024). With thoughtful, ethically informed use, however, it may support more inclusive and responsive language-learning environments.

Future research should test EALEM in diverse educational contexts, including multilingual classrooms and under-resourced settings, in order to examine its feasibility, limitations, and pedagogical impact. Such research should move beyond the question of what AI can do and address the more demanding question of what AI should do in education (Dwivedi et al., 2023). The future of GenAI in language education will depend less on the technology itself than on the values, frameworks, and practices that shape its use. If approached critically and ethically, GenAI may not become a substitute for human teaching but a tool that supports educational justice, linguistic diversity, and reflective learning.

References

- Barnett, S. (2025, August 18). Teachers are trying to make AI work for them. *Wired*. <https://www.wired.com/story/teachers-using-ai-schools>
- Canagarajah, S. (2012). *Translingual Practice: Global Englishes and Cosmopolitan Relations* (1st ed.). Routledge. <https://doi.org/10.4324/9780203073889>

- Chaudhry, A., Cukurova, M., & Luckin, R. (2022). *A transparency index framework for AI in education*. arXiv. <https://arxiv.org/abs/2206.03220>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., ... & Wright, R. (2023). Opinion Paper: “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges, and implications of generative conversational AI for research, practice, and policy. *International journal of information management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Ellis, R. (2003). *Task-based language learning and teaching*. Oxford University Press.
- Gabriel, S. (2024). Generative AI and educational (in)equity. *Proceedings of the International Conference on AI Research*, 4(1), 133-142. <https://doi.org/10.34190/icair.4.1.3153>.
- García-López, I. M., Sánchez-Vera, M. M., Solano-Fernández, I. M., & Pérez-Hernández, D. (2025). Ethical and regulatory challenges of generative AI in education. *Frontiers in Education*, 10, 1565938. <https://doi.org/10.3389/feduc.2025.1565938>
- Giannakos, M. (2024). The promise and challenges of generative AI in education. *Behaviour & Information Technology*, 43(11), 1299–1312. <https://doi.org/10.1080/0144929X.2024.2394886>
- Krashen, S. D. (1982). *Principles and practice in second language acquisition*. Pergamon Press.
- Lantolf, J. P., & Thorne, S. L. (2006). *Sociocultural theory and the genesis of second language development*. Oxford University Press.
- Lee, S., Choe, H., Zou, D., & Jeon, J. (2026). Generative AI (GenAI) in the language classroom: A systematic review. *Interactive Learning Environments*, 34(1), 335–359. <https://doi.org/10.1080/10494820.2025.2498537>
- Li, B. (2025). Two years of innovation: A systematic review of empirical generative AI research in language learning and teaching. *Computers & Education: Artificial Intelligence*, 8, 100445. <https://doi.org/10.1016/j.caeai.2025.100445>
- Moorhouse, B. L. (2025). *Generative artificial intelligence and language teaching*. Cambridge University Press. <https://doi.org/10.1017/9781009618823>
- Nyaaba, M., Wright, A., & Choi, G. L. (2024). *Generative AI and digital neocolonialism in global education*. <https://doi.org/10.48550/arXiv.2406.0296>
- OECD. (2023). *AI in education: Ensuring equity and transparency*. OECD Publishing. <https://doi.org/10.1787/ai-education-en>
- Pennycook, A. (2017). *Posthumanist applied linguistics*. Routledge.
- Richards, J. C., & Rodgers, T. S. (2014). *Approaches and methods in language teaching* (3rd ed.). Cambridge University Press.
- Stokes, K. (2025, September 2). Minnesota schools embrace AI—cautiously. *Axios Twin Cities*. <https://www.axios.com/local/twin-cities/2025/09/02/ai-back-to-school-minnesota>
- Swain, M. (1995). Three functions of output in second language learning. In G. Cook & B. Seidlhofer (Eds.), *Principle and practice in applied linguistics* (pp. 125–144). Oxford University Press.
- U.S. Department of Education. (2023). *Artificial intelligence and the future of teaching and learning*. <https://files.eric.ed.gov/fulltext/ED631097.pdf>
- UNESCO. (2023). *Guidance for generative AI in education and research*. UNESCO Publishing. <https://unesdoc.unesco.org/ark:/48223/pf0000386693>
- Vygotsky, L. S. (1978). *Mind in society: The development of higher psychological processes*. Harvard University Press.
- Weng, Z., & Fu, Y. (2025). Generative AI in language education: Bridging divide and fostering inclusivity. *International Journal of Technology in Education*, 8(2), 395-420. <https://doi.org/10.46328/ijte.1056>

Yan, Y., Liu, H., & Chau, T. (2025). A Systematic Review of AI Ethics in Education: Challenges, Policy Gaps, and Future Directions. *Journal of Global Information Management (JGIM)*, 33(1), 1-50. <https://doi.org/10.4018/JGIM.386381>

Eirini Ioanna Delmadorou (reniadelmar@gmail.com) is a philologist, English language educator, and certified adult trainer. She holds a BA in Philosophy from the National and Kapodistrian University of Athens and an MA in Teaching English as a Foreign/International Language from the Hellenic Open University. Her academic background includes specialisation in Political and Economic Philosophy, Bioethics, and the intersection of Philosophy and Technology. She has extensive teaching experience in secondary education and private tutoring, teaching Ancient and Modern Greek, History, Philosophy, Latin, and English as a Foreign Language. She has also completed teaching placements at the Model High School of the Ionideios School of Piraeus.