



Research Papers in Language Teaching and Learning

Vol. 16, No. 1, March 2026, 31-46

ISSN: 1792-1244

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

This article is issued under the [Creative Commons License Deed. Attribution 3.0 Unported \(CC BY 3.0\)](#)

From Red Ink to Algorithm: Reimagining Feedback Cultures with GenAI in EFL Writing

Ioanna Nifli

Emerging debates in applied linguistics and educational technology have increasingly turned to the transformative role of Generative AI, and more specifically, LLM-powered tools such as ChatGPT, Claude, or Gemini, in reshaping feedback practices in second language writing. Beyond language education, these debates reflect a broader pedagogical reckoning with long-standing behaviourist feedback traditions that prioritise error detection, compliance, and correction over meaning-making and learner agency. Situated within this evolving landscape, the present theoretical inquiry examines how this emerging technology intersects with entrenched feedback cultures in EFL classrooms. The “red-pen” logic of feedback is conceptualised as part of a wider instructional paradigm that frames learning as error elimination through external reinforcement, an orientation historically aligned with behaviourist models of correction, compliance, and measurable accuracy. Drawing on sociocultural theory, feedback literacy, and affect-informed pedagogy, the article critiques the prevailing model of feedback as unidirectional error correction and argues for a shift toward dialogic, co-constructed feedback cultures mediated by both human and algorithmic actors. To account for the deep-rooted patterns that shape learners’ orientations toward feedback, the paper introduces the notion of Affective-Epistemic Feedback Habitus—a dispositional framework formed through repeated exposure to authoritative correction, behaviourist assessment logics, and high-stakes evaluation. This habitus is conceptualised as a durable set of emotional and epistemological “defaults” that govern what feedback feels like (e.g., threat, reassurance, shame, safety), what it counts as (e.g., verdict vs. invitation), and how learners are oriented to act upon it (e.g., comply, conceal, negotiate, revise). It shapes how students emotionally and cognitively interpret LLM-mediated feedback, often predisposing them to perceive algorithmic suggestions as final judgments rather than as opportunities for reflection or revision. By exploring the pedagogical affordances and risks of LLM-generated feedback in this context, the paper argues that these tools offer both promise and limitations. While their non-judgmental tone and real-time support can scaffold learner autonomy and metacognitive engagement, their uptake within correction-oriented pedagogical cultures risks reproducing epistemic dependency, instrumental revision practices, and motivational detachment. The article advocates repositioning LLMs not as autonomous evaluators but as pedagogical interlocutors embedded in culturally responsive, ethically mediated writing instruction. Ultimately, it reframes feedback not as a static transmission of correctness, but as a situated, emotionally resonant process co-constructed within specific educational ecologies.

Keywords: Generative AI; Large Language Models; EFL Writing Feedback; Feedback Cultures; Behaviourism; Learner Agency; Feedback Literacy; Dialogic Pedagogy; Sociocultural Theory; Affective Feedback; Teacher Mediation

1. Introduction

The rise of Generative Artificial Intelligence (GenAI) in education has reshaped how learners engage with language, writing, and feedback. LLM-powered tools such as ChatGPT, Claude, and Gemini now deliver instant, personalised responses, from grammar correction to idea expansion and stylistic refinement, promising to democratise access to feedback, enhance writing fluency, and support metacognitive growth (Cavaleri et al., 2024:311; Liu & Wang, 2023: 88). Yet this promise is neither neutral nor universally realised. At a systemic level, the emergence of GenAI intersects with a longstanding pedagogical tension across educational contexts between behaviourist models of learning centred on stimulus, correction, and reinforcement, and sociocultural, dialogic approaches that conceptualise learning as meaning-making, participation, and agency.

Historically, the “red pen” has functioned as both a material and symbolic instrument of correction-oriented pedagogy, marking deviation, signalling error, and enforcing compliance (Semke, 1984:195). Across diverse educational systems and subject domains, correction-heavy feedback practices have been shown to undermine learner confidence, narrow revision to surface-level edits, and reduce writing to error avoidance rather than rhetorical development. Research consistently links such practices to heightened anxiety, diminished motivation, and shallow learning outcomes (Horwitz et al., 1986:125; Ryan & Henderson, 2018:884). The red-pen logic thus reflects a broader instructional paradigm in which feedback operates primarily as external regulation rather than as mediated support for learning.

While global debates have increasingly examined the pedagogical and ethical implications of LLM integration, the effects of these technologies remain uneven across educational systems shaped by distinct linguistic, cultural, and assessment traditions. In many EFL contexts—particularly those characterised by high-stakes assessment, strong teacher authority, and accuracy-oriented curricula—feedback practices continue to be dominated by correctional norms and emotionally charged evaluative dynamics. The red-pen model examined in this paper, therefore, reflects a widely recognisable educational habit of positioning feedback as behavioural control: detecting deviation, marking error, and rewarding conformity. When feedback functions primarily as external regulation, learner attention is systematically narrowed to error avoidance rather than to the development of voice, meaning, and rhetorical agency.

In exam-oriented EFL systems, writing pedagogy has frequently prioritised grammatical accuracy and lexical appropriateness over rhetorical development or authorial voice (Lyriqkou, 2021). Feedback in such contexts is often delivered through visibly corrective, hierarchical, and emotionally consequential practices that many learners experience as judgment rather than an invitation to revise, fostering anxiety, demotivation, and superficial revision strategies. Within these instructional ecologies, LLM-generated feedback holds the potential both to disrupt established norms and to amplify them.

Although GenAI tools can prompt more frequent revision and quicker feedback, their pedagogical impact is not inherent but contingent on how they are framed, mediated, and taken up. In correction-oriented

feedback cultures, learners may default to using GenAI as an automated grammar checker, while teachers may tacitly legitimise this use by aligning AI feedback with error elimination rather than metacognitive engagement. Under these conditions, GenAI risks functioning as a “digital red pen”: efficient, immediate, and scalable, yet pedagogically reductive. This risk does not arise because GenAI is intrinsically correctional, but because existing feedback scripts, shaped by behaviourist legacies, assessment incentives, and authority-driven norms, can domesticate a dialogic technology into an instrumental workflow. In high-stakes, exam-driven environments, the path of least resistance is often instrumental use rather than dialogic meaning-making.

This paper, therefore, treats LLM adoption as a grand challenge of feedback culture design across EFL contexts: whether new technologies will be used to extend corrective regimes, accelerating error detection and compliance, or to cultivate dialogic, learner-centred revision practices that redistribute interpretive agency.

Accordingly, the paper critically examines how LLM-powered tools interact with feedback cultures that is, the normative and affective patterns shaping how feedback is given, received, and interpreted in localised pedagogical settings (Carless & Winstone, 2020:153). Drawing on sociocultural theory, feedback literacy, and affect-informed pedagogy, it argues that feedback is not merely technical text improvement but a relational, culturally embedded act reflecting broader pedagogical ideologies. Particular attention is paid to how behaviourist feedback traditions may be inadvertently reproduced through GenAI when critical AI literacies and teacher mediation are absent. The analysis explores how LLMs may either reinforce hierarchical, judgmental paradigms or foster dialogic, agency-oriented models of feedback, positioning them not as autonomous correctors but as scaffolds for learner authorship, metacognition, and emotional safety when mediated with pedagogical and cultural sensitivity.

2. Theoretical Framework: Feedback Cultures and Sociocultural Perspectives

2.1 Feedback Cultures

In recent years, the concept of feedback cultures has emerged as a critical lens for understanding how feedback is shaped not only by pedagogical intentions but also by institutional, relational, and emotional norms. Carless and Winstone (2020:156) define feedback cultures as ‘the shared understandings, expectations and values that underpin feedback practices and interactions in a given learning context’. Rather than viewing feedback as a neutral, transactional act, this perspective highlights its social and cultural construction, rooted in power relations and institutional histories. This paper additionally situates correction-dominant feedback cultures within a wider behaviourist orientation to teaching, where feedback functions primarily as external reinforcement and error suppression, rather than as mediation for meaning-making and development.

In EFL writing education, feedback cultures are frequently characterised by correction-oriented paradigms and strong teacher authority. Ajjawi and Boud (2017: 252–260) observe that feedback is often experienced as asymmetrical, with teachers positioned as gatekeepers of correctness and learners expected to decode, accept, and comply. Such cultures privilege surface-level linguistic accuracy over rhetorical development, metacognitive awareness, and learner voice. While these dynamics are particularly visible in exam-oriented EFL contexts, they are by no means context-specific. Across many compliance-driven and assessment-intensive educational systems, feedback is institutionalised as surveillance of error rather than as dialogue about meaning-making, communicative intent, or writer identity. In these settings, visible correction practices—often symbolised by the “red pen”—come to

signify not only linguistic deviation but also academic judgement, reinforcing hierarchical feedback relations and limiting opportunities for dialogic engagement (Vlanti, 2012: 92).

Feedback cultures are also affectively charged as students interpret feedback through emotional filters shaped by prior experiences, teacher relationships, and classroom hierarchies. In high-stakes environments, such as university entrance exams or international certification tests, feedback can provoke anxiety, shame, or disengagement. Yet, despite this emotional impact, the affective dimension remains marginalised in both research and practice (Ryan & Henderson, 2018:884).

To address this gap, this paper introduces the concept of *Affective-Epistemic Feedback Habitus*: the internalised emotional and epistemological dispositions learners acquire through repeated exposure to culturally dominant feedback practices. Drawing on Bourdieu's (1977) theory of habitus, affect-informed pedagogy (Ushioda, 2011; Dewaele, 2020), and sociocultural feedback research (Carless & Boud, 2018; Ajjawi & Boud, 2017a), this concept captures how feedback is felt, trusted, and acted upon. Unlike feedback cultures, which describe external norms, or scripts, which outline expected routines, this habitus reflects a deeply embodied orientation shaping learners' intuitive reactions, compliance, resistance, anxiety, or detachment, especially in systems where correction, authority, and emotional vulnerability are intertwined.

More specifically, the *Affective-Epistemic Feedback Habitus* refers to a patterned "readiness" to interpret feedback through (a) affective expectancy (anticipating criticism, threat, or relief), and (b) epistemic positioning (treating the feedback source as unquestionable authority vs. negotiable interlocutor). In SLA terms, it helps explain why two learners receiving similar feedback may diverge sharply in uptake: one may revise strategically and reflectively, while another may either comply mechanically or disengage, depending on prior socialisation into what feedback *is* **and what it does to** the self. In EFL writing, where identity, voice, and public evaluation intersect, this habitus can amplify foreign language anxiety (Horwitz et al., 1986:125), shape motivation and self-concept (Ushioda, 2011:199), and condition whether feedback becomes "usable information" or an affective threat requiring avoidance (Ryan & Henderson, 2018: 884). The construct also clarifies a crucial pedagogical point: feedback literacy is not developed in a vacuum; it is mediated by dispositions. Where learners have internalised correction-as-judgment, they may possess procedural knowledge of "fixing errors" yet lack the epistemic confidence and emotional safety needed to negotiate feedback dialogically (Carless & Boud, 2018: 1315).

2.2 Sociocultural and Affective Pedagogy

To theorise feedback beyond static delivery, this paper draws on sociocultural learning theory, particularly Vygotsky's concepts of mediation (1978) and the Zone of Proximal Development (ZPD). From this perspective, feedback is most effective when it functions as a mediational tool, scaffolding learners' movement from current to potential performance within a supportive social context (Lantolf & Thorne, 2006:79). Unlike mechanistic correction, feedback is conceived as a dialogic process co-constructed between teacher and learner, embedded in the dynamic interplay of cognition, emotion, and social interaction.

The integration of LLM-based tools into feedback processes both aligns with and challenges this model. On the one hand, GenAI can provide timely, personalised suggestions that scaffold independent revision; on the other, its lack of relational and emotional attunement risks flattening the dialogic space, reducing feedback to surface-level edits detached from learner identity and intention. This raises critical questions about the locus of mediation, whether it resides in the algorithm, the teacher, or their interaction, and about how such mediation resonates with learners' affective and cultural expectations. It also foregrounds pragmatic concerns: culture-sensitive feedback depends not only on *what* is suggested, but

on *how* it is said, what it presupposes, and how learners interpret it as a speech act involving stance, politeness, and implied authority.

Affective dimensions—long marginalised in feedback research—are central to understanding learner engagement in EFL writing, where personal expression, linguistic vulnerability, and public evaluation intersect. Motivation and emotion shape learners’ beliefs, self-concept, and agency (Ushioda, 2011: 199), while affective responses to feedback, from pride to shame, directly influence whether learners revise, persist, or disengage (Dewaele, 2020). In many EFL contexts characterised by hierarchical teacher–student relations, strong evaluative traditions and authority-driven feedback practices, teacher commentary often carries moral and identity-related weight, positioning feedback as judgment rather than guidance. In such settings, replacing human feedback with algorithmic output may weaken the relational and motivational bonds that sustain learner effort, particularly when feedback has historically functioned as a source of validation, reassurance, or control. More broadly, this concern extends to any instructional context in which feedback is treated primarily as “control information” rather than as “learning dialogue”: affective safety is not a peripheral consideration but a necessary condition for sustained revision, risk-taking, and the development of authorial voice in second language writing.

To address these tensions, the paper adopts the lens of feedback literacy, understood as the capacity to understand, evaluate, and use feedback productively (Carless & Boud, 2018:1318). Feedback literacy extends beyond decoding comments to include interpreting purpose, negotiating meaning, and integrating feedback into one’s learning trajectory. In AI-mediated contexts, this literacy must be explicitly cultivated: learners require guidance not only in using GenAI tools, but in critically engaging with their suggestions. Accordingly, the paper treats AI literacies as a complementary set of competencies, critical, ethical, and interactional, needed to interpret LLM outputs, recognise their limitations, and maintain authorship. Because GenAI introduces a new epistemic actor into the feedback ecology, feedback literacy must expand to include understanding probabilistic outputs and evaluating suggestions against rhetorical intent and task criteria, rather than accepting them as authoritative correction.

In AI-augmented feedback ecologies, agency must be preserved and pedagogically scaffolded. Learners should be positioned as active participants in a dialogic feedback loop, with AI as support rather than a surrogate. Teachers, in turn, act as co-mediators, helping students interpret, critique, and personalise LLM-generated feedback in ways that are both culturally and emotionally responsive. Crucially, this co-mediation is not only “pedagogical” but also “pragmatic”: when LLMs produce feedback in ways that misalign with local norms of politeness, directness, stance, or implicature, learners may misread the intent or weight of suggestions. Therefore, culture-sensitive feedback design must consider the pragmatics of LLM interaction, not only its grammatical correctness.

2.3 Correction-Dominant Feedback Cultures in EFL Writing: A Global Problem with Local Manifestations

Correction-focused feedback practices remain deeply entrenched in EFL writing instruction across many educational systems, reflecting broader institutional, ideological, and assessment-driven traditions. From early schooling onward, learners in numerous exam-oriented contexts are socialised into feedback cultures where writing is evaluated primarily through annotation, error marking, and visible correction rather than dialogic engagement with meaning, voice, or rhetorical intent. These practices are not unique to any single national context; rather, they represent a widely recognisable legacy of behaviourist pedagogy in which learning is operationalised as error reduction through external reinforcement, a logic shown to constrain deep learning, undermine emotional well-being, and limit sustained writing development when it dominates feedback practices (Semke, 1984: 195; Ryan & Henderson, 2018: 884).

Within this broader landscape, EFL writing classrooms illustrate how high-stakes assessment regimes concretely shape feedback practices. In many educational systems, exam-driven curricula have historically framed writing as a test of grammatical accuracy and lexical control rather than communicative effectiveness or rhetorical development. Such regimes, characterised by standardised criteria, time pressure, and accountability structures tend to reward formulaic responses and surface-level correctness over exploratory drafting, voice, or audience awareness. Consequently, writing instruction often becomes a high-risk, performance-oriented activity in which success is defined by error minimisation rather than the cultivation of authorial agency. While the intensity of these pressures varies across contexts, their pedagogical consequences are consistently observable in EFL classrooms governed by standardisation, external evaluation, and outcome-driven assessment logics.

As a result, feedback is frequently framed as evaluative control rather than as a learning dialogue. Teachers—operating under institutional and societal pressure to deliver measurable outcomes—are positioned as grammatical arbiters whose role is to detect deviation from standard usage, mark errors (often in red pen), and return work with limited explanation or scaffolding. The red ink itself functions as a powerful pedagogical symbol, signalling judgment, finality, and correctness. Students are rarely invited into reflective dialogue about their rhetorical choices; instead, feedback is experienced as a verdict to be accepted and implemented rather than negotiated or questioned—if revision is encouraged at all (Semke, 1984: 195).

Such correction-dominant practices carry substantial affective weight. Writing becomes a high-risk act, exposing learners to correction, disappointment, and potential loss of face. Rather than cultivating revision as a recursive, exploratory process, feedback often reinforces surface-level edits and compliance with perceived norms. Over time, learners may become rhetorically risk-averse, overly dependent on memorised templates and model answers, or reluctant to develop an individual voice. The emotional cost is particularly acute for students who struggle with grammatical accuracy or who internalise early labels of being “weak” in English—labels closely associated with foreign language anxiety and diminished self-efficacy (Horwitz et al., 1986: 125).

Through repeated exposure to such feedback regimes, these tendencies gradually become dispositional rather than situational. Learners develop what this paper conceptualises as an *Affective-Epistemic Feedback Habitus*, shaped by authoritative correction, high-stakes evaluation, and emotionally charged feedback encounters. This habitus predisposes learners to interpret feedback as final, to internalise correction as personal failure, and to avoid the reflective, iterative dimensions of revision. It functions as an interpretive filter through which all feedback, whether teacher-delivered or GenAI-generated, is processed. More specifically, it normalises compliance over inquiry (epistemic stance) and threat over possibility (affective stance), making “fixing errors” feel like the only legitimate response to feedback, even when meaning-focused or rhetorical revision would be pedagogically preferable. In this sense, the red pen operates not merely as a tool but as a socialising practice: it trains learners to locate authority outside the self and to treat correctness as externally imposed truth rather than jointly constructed meaning.

The emotional climate surrounding feedback is rarely addressed explicitly in classrooms. Across many EFL contexts, teachers receive limited training in affect-sensitive pedagogy, and institutional support for formative, dialogic feedback remains constrained by curricular demands. Societal and parental expectations for academic success further reinforce accuracy-oriented norms, sustaining a pervasive climate of anxiety around writing. While such dynamics are not exclusive to EFL, they are intensified by the symbolic power of English as a gatekeeping language for educational mobility and professional opportunity (Vogt & Tsagari, 2014: 53).

The lived consequences of these feedback cultures are reflected in familiar classroom experiences. A student receives an essay densely marked with red annotations, misspellings, verb tense corrections, punctuation errors—accompanied by a grade and a single directive (“Revise”), with no comment on argument or organisation. Another learner, preparing for an international language examination, is instructed to memorise fixed essay templates and discouraged from rhetorical experimentation, receiving the implicit message that conformity, not expression, ensures success. Such vignettes, common across exam-driven EFL systems, illustrate how institutional pressure, pedagogical intent, and emotional impact converge in correction-dominant feedback cultures.

Within these contexts, GenAI-enhanced technologies introduce both risk and possibility. If learners approach GenAI primarily as a grammar checker, and if teachers implicitly legitimise this use by aligning AI feedback with error elimination rather than metacognitive engagement, GenAI risks becoming a digital red pen—efficient, immediate, and scalable, yet pedagogically reductive. This risk is not hypothetical. In accuracy-driven systems, the most socially rewarded and time-efficient use of GenAI is often surface correction (“fix my grammar,” “improve my vocabulary”), because it aligns seamlessly with assessment incentives, parental expectations, and entrenched classroom routines (Vogt & Tsagari, 2014: 57).

Crucially, platform affordances can amplify this trajectory. The ease of copy–paste revision, one-click rewriting, and “polish” prompts encourages workflows where texts are edited into acceptability rather than revised into meaning, and improvement is equated with linguistic cleanup rather than conceptual rethinking. Over time, this can normalise a revision habitus of outsourcing, in which learners implement changes they cannot explain, weakening feedback literacy and diminishing the inner dialogue that makes feedback educative (Nicol, 2021:24; Carless & Boud, 2018:1318). In such ecologies, a dialogic technology is domesticated into an instrumental workflow.

Yet disruption also creates space for reconfiguration. With teacher-led mediation, explicit attention to affect and pragmatics, and the cultivation of feedback and AI literacies, LLM systems could support a shift toward process-oriented writing and dialogic revision practices that redistribute interpretive agency. Achieving this shift requires more than adopting new tools; it demands critical reflection on the historical, emotional, and institutional architectures of feedback that shape how learners interpret and respond to feedback, whether human or algorithmic. It also requires shared understandings of what constitutes meaningful AI use in revision (e.g., questioning, justifying, resisting over-editing), alongside sustained teacher professional development addressing pedagogical orchestration, ethics, and the pragmatics of human–LLM interaction.

2. LLMs and Feedback Mediation

3.1 Affordances of GenAI Feedback Tools

LLM-based tools such as ChatGPT, Claude, Gemini, and QuillBot have introduced new paradigms of writing support by delivering real-time, personalised, context-aware feedback that is typically non-judgmental in tone. Unlike traditional feedback, which is often delayed, constrained by teacher workload, and emotionally charged, these systems can address concerns ranging from grammatical accuracy to rhetorical structure and lexical sophistication (Cavaleri et al., 2024: 2; Liu & Wang, 2023: 108).

A key pedagogical affordance of LLM-mediated feedback lies in its capacity to scaffold iterative revision. Through successive prompts and paraphrasing cycles, learners can refine their writing via experimentation, reflection and revision, aligning with Vygotsky’s Zone of Proximal Development, where

learning is optimally supported just beyond independent capability (Lantolf & Thorne, 2006: 80). Used intentionally, GenAI can function as a digital scaffold, enabling learners to identify gaps, explore alternatives and reflect on linguistic choices without immediate penalty.

The dialogic capabilities of advanced GenAI platforms, especially those with memory and adaptive feedback, further enable metacognitive engagement. Learners can ask why a sentence is awkward, how to strengthen an argument, or whether a word choice suits a formal register. Such “explain-this-feedback” interactions promote conceptual awareness of genre, register, and audience, key aspects of academic writing literacy, and foster what Nicol (2021: 24) terms *feedback as an inner dialogue*, where meaning is co-constructed through formative engagement.

In this light, GenAI can reposition feedback from a static, teacher-delivered product to a dynamic, learner-initiated process. When framed pedagogically, these tools have the potential to empower students as proactive agents in their writing development, using AI not as a shortcut but as a dialogic partner in composition. This dialogic potential is also pragmatic: LLMs can model stance, hedging, politeness strategies, and audience-sensitive reformulation, dimensions that matter for culture-sensitive feedback and for learners’ perception of whether feedback is respectful, supportive, and actionable.

3.2 Tensions Between GenAI Feedback and Correction-Dominant Feedback Norms in EFL Contexts

Despite their dialogic affordances, the integration of LLM-mediated feedback systems into EFL writing instruction generates significant pedagogical and cultural tensions across diverse educational contexts. Central among these is a recurrent mismatch between GenAI’s probabilistic, non-authoritative feedback style and entrenched correction-dominant feedback norms that characterise many exam-oriented and accuracy-driven language education systems worldwide. In such contexts, learners are often socialised to interpret feedback as a final evaluative judgment, closely associated with grades, visible correction, and error identification, rather than as an open-ended invitation to revise, reflect, or negotiate meaning.

This tension is particularly visible in EFL settings shaped by behaviourist legacies, where writing instruction has historically privileged linguistic correctness, rule adherence, and conformity to standardised norms. Within these systems, feedback is expected to provide unequivocal answers to the binary question of correctness. Consequently, GenAI’s non-directive suggestions (e.g., “*You may wish to consider revising this sentence for clarity*”) may be perceived as insufficiently authoritative or pragmatically ambiguous. Empirical observations from Southern European and other exam-oriented contexts indicate that learners frequently seek categorical confirmation when encountering AI feedback, asking variations of “Is this right or wrong?” (Garcia & Lee, 2023: 89). This response signals a deeper epistemological tension: GenAI’s context-sensitive, probabilistic recommendations contrast sharply with the rule-based certainties cultivated by traditional EFL feedback regimes.

In many EFL contexts internationally, these tensions are particularly salient. Learners are often socialised within highly evaluative feedback cultures in which teacher authority, visible correction practices (e.g., red-pen marking), and high-stakes assessment structures dominate writing instruction. For learners shaped by an *Affective-Epistemic Feedback Habitus* that equates feedback with authoritative correctness, GenAI’s mitigated, probabilistic guidance may be experienced as confusing, insufficiently decisive, or even unreliable unless explicitly mediated by teachers. Importantly, these dynamics are not confined to a single national setting. Comparable patterns have been documented across diverse EFL contexts where assessment incentives, institutional accountability, and parental expectations prioritise surface accuracy

and error elimination over rhetorical development, authorial agency, and dialogic engagement with feedback.

Recent scholarship on the pragmatics of LLM-mediated interaction further complicates assumptions that AI-generated feedback is neutral or universally interpretable. Research in computational pragmatics and AI discourse analysis demonstrates that LLMs systematically encode politeness strategies, epistemic hedging, indirectness, and face-preserving formulations rooted primarily in Anglophone academic discourse norms (Brown & Levinson, 1987: 66–68; Danescu-Niculescu-Mizil et al., 2013: 8; Hovy & Spruit, 2016: 592). While such pragmatic features may soften affective impact, they also introduce interactional ambiguity in educational cultures accustomed to directive, authority-driven feedback. LLM-generated feedback typically privileges mitigated suggestions over categorical correction, reflecting probabilistic alignment rather than pedagogical intentionality (Bender et al., 2021: 615; Floridi et al., 2018: 694). In many EFL contexts, this pragmatic mismatch can lead learners either to dismiss AI feedback as indecisive or to overinterpret it as covertly authoritative, thereby reinforcing compliance rather than dialogic engagement. These dynamics underscore the need to conceptualise GenAI feedback not only as a cognitive or affective intervention, but as a socio-pragmatic act whose interpretation is culturally mediated and pedagogically consequential.

These tensions are further intensified by what may be described as feedback shortcutting. Learners may rely on tools such as Grammarly or ChatGPT to eliminate surface errors while bypassing the cognitive work of developing writing strategies. In assessment-driven EFL systems, where correctness is rewarded more visibly than rhetorical or metacognitive growth, this reliance risks reducing feedback to linguistic outsourcing (Vogt & Tsagari, 2014:57). This mechanism explains how an intuitively dialogic technology can become a digital red pen: dominant success metrics (error-free text) steer usage toward the fastest visible gains (surface correction), reinforcing compliance-based revision habits and weakening reflective uptake. Moreover, a subtle standardisation effect emerges, as LLMs, optimised for fluent, norm-aligned prose, implicitly privilege conventional phrasing and generic correctness, especially when learners request *“better”* or *“more advanced”* English without articulating rhetorical intent. Over time, this norm-alignment risks reproducing correctional hierarchies under a new technological interface, positioning learners as implementers of external edits rather than negotiators of meaning (Ajjawi & Boud, 2017b: 252–260; Carless & Winstone, 2020:15 6).

This risk is amplified by techno-solutionist discourses that frame GenAI as an efficiency-enhancing remedy for long-standing feedback challenges such as teacher workload, delayed correction, or error density. When adopted primarily as a tool for speed, scalability, or cost reduction, GenAI may automate and intensify existing correctional paradigms rather than transform them. In EFL systems shaped by high-stakes assessment, techno-solutionist adoption often positions GenAI as a mechanism for faster error detection and surface-level optimisation, thereby reinforcing behaviourist feedback logics under the guise of innovation. Without deliberate pedagogical reframing, GenAI does not disrupt the red-pen culture; it digitises it.

Teacher roles are also implicated in this shift. If LLM-powered feedback is positioned as a substitute for teacher input, educators risk marginalisation, with potential erosion of both pedagogical authority and the relational trust that sustains learner motivation (Semke, 1984:197). In the absence of structured mediation, LLM-generated feedback may operate in a pedagogical vacuum, technically proficient yet affectively and culturally misaligned. This misalignment is not only emotional but pragmatic: LLM feedback may mishandle politeness, directness, implicature, humour, or culturally preferred degrees of explicitness, leading learners to misinterpret suggestions as overly authoritative, insufficiently clear, or excessively interventionist. Culture-sensitive integration, therefore, requires sustained attention to the

pragmatics of human–LLM interaction, including how stance, interpersonal alignment, and evaluative force are communicated through AI feedback.

Addressing these tensions necessitates embedding LLM-mediated feedback within teacher-facilitated, dialogic feedback cultures. Learners require explicit support in developing critical feedback literacy: the capacity to interpret, critique, adapt, and, when appropriate, reject algorithmic suggestions in line with rhetorical purpose, audience awareness, and communicative intent. Such discernment ensures that GenAI functions as a dialogic support rather than a decontextualised corrector, preserving both pedagogical integrity and learner agency. Crucially, this training must be framed as part of AI literacies alongside feedback literacy, encompassing interpretive competencies (evaluating output quality), interactional competencies (formulating productive prompts), and ethical competencies (maintaining authorship and academic integrity). In practice, this involves designing feedback tasks that make *thinking with feedback* visible, requiring justification, annotation, and comparison, so that the learning outcome is not merely a cleaner text, but a more agentive, reflective writer. Taken together, these dynamics illustrate how GenAI’s pedagogical impact is not determined by its dialogic capacity alone, but by the feedback cultures into which it is embedded.

3.3 Emotional and Motivational Implications

One of the most under-theorised yet critical dimensions of GenAI feedback is its emotional impact on learners. In feedback cultures where fear of judgment, performance anxiety, and correction fatigue are prevalent, LLM-mediated feedback offers a potential affective shift. By presenting suggestions in a neutral, impersonal tone, GenAI may reduce the emotional sting often associated with teacher correction. For students who link teacher feedback with failure or embarrassment, AI can provide a psychologically safer space for experimentation and risk-taking in writing (Ryan & Henderson, 2018: 884). Emotional safety in this context is not trivial: it underpins creativity, cognitive engagement, and willingness to take linguistic risks, particularly in writing, where personal expression meets linguistic vulnerability.

When framed pedagogically, GenAI can help learners move beyond formulaic structures toward more autonomous and confident expression. This aligns with Ushioda’s (2011:202) view of motivation as a relational and affective construct shaped by the learner’s emotional experience of instructional activities and evaluative interactions.

Yet the absence of affectively salient human feedback carries risks. *Authorship blurring*—the uncritical incorporation of AI suggestions—may dilute ownership of the text, complicating questions of originality and reducing opportunities for metacognitive reflection. Learners may revise “correctly” without understanding why, resulting in superficial gains. Overreliance on GenAI can also foster motivational detachment: if revision is routinely outsourced, writing may be perceived less as personal expression or growth and more as a mechanical process to be optimised algorithmically. This instrumentalisation of writing risks weakening intrinsic motivation and eroding confidence in one’s authorial voice.

To mitigate these risks, GenAI must be positioned not as an evaluator but as a conversational partner whose feedback is to be interpreted, adapted, or sometimes rejected. Teachers should help learners develop discernment, ownership, and critical trust in both human and AI feedback. Without such mediation, the promise of transforming feedback into a dialogic, learner-centred process may remain unrealised, particularly in cultures where feedback is historically bound to emotional vulnerability and academic authority. From a pragmatic perspective, this also requires attention to how LLM feedback performs stance: hedges (“perhaps,” “consider”), directives (“change X to Y”), and evaluative language can be interpreted differently across cultures and proficiency levels, affecting whether learners perceive feedback as supportive, ambiguous, or authoritative.

4. Reimagining Feedback Cultures: Pedagogical Implications

In many EFL contexts, these tensions become particularly visible. Learners are often socialised within highly evaluative feedback cultures dominated by teacher authority, red-pen correction, and high-stakes assessment. For learners shaped by an *Affective-Epistemic Feedback Habitus* that equates feedback with authoritative correctness, GenAI's mitigated or non-directive guidance may prompt confusion, mistrust, or dismissal unless explicitly mediated by teachers. Comparable dynamics have been documented across EFL settings where institutional expectations, assessment incentives, and parental pressures prioritise surface accuracy over rhetorical development and dialogic engagement.

Rather than treating LLMs as autonomous evaluators, educators should embed them within frameworks that protect learner agency, emotional safety, and cultural specificity. Unmediated use risks automating existing correctional paradigms; thoughtful integration can instead position LLMs as dynamic mediators of writing development, prompting reflection, negotiation, and ownership. Achieving this transformation hinges on two interdependent priorities: repositioning teachers as feedback mediators and cultivating learner feedback literacy in LLM-rich environments. A third, cross-cutting priority is the explicit development of AI literacies and sustained teacher professional development, ensuring that AI does not function as an invisible curriculum silently reshaping what counts as writing, revision, and improvement.

4.1 Pedagogical Repositioning of Teachers

The introduction of GenAI into feedback processes necessitates a shift in the teacher's role—from evaluator to feedback mediator. Teachers are called to move beyond acting as final arbiters of correctness and instead function as interpreters and co-constructors of meaning, guiding learners through GenAI outputs with critical, affective, and contextual awareness (Winstone & Carless, 2020: 67).

This repositioning does not diminish teacher expertise; rather, it amplifies it by foregrounding the orchestration of meaningful interaction between learner and machine. Effective mediation involves helping students understand not only what GenAI suggests, but why, how, and when those suggestions align, or conflict, with communicative goals, genre conventions, and audience expectations, particularly when AI feedback is vague, formulaic, or culturally misaligned. In this way, teachers preserve feedback as dialogue, ensuring that AI functions as a catalyst for reflection rather than a decontextualised corrector (Nicol, 2021:27). Crucially, such mediation also supports pragmatic alignment, enabling learners to interpret stance, hedging, directness, and politeness in AI feedback and to negotiate revisions that remain rhetorically and culturally appropriate.

Teachers must also attend to the emotional dimension of AI-mediated feedback, supporting learners who may experience confusion, dependency, or diminished authorship when working with algorithmic suggestions. This affective scaffolding is essential for sustaining motivation and trust—elements often absent from purely algorithmic exchanges (Ryan & Henderson, 2018:885). Teacher professional development is therefore central: educators require not only technical familiarity with GenAI tools, but also AI literacies, critical pedagogical frameworks, and guidance in orchestrating dialogic human–AI–learner feedback interactions.

4.2 Cultivating Feedback Literacy in AI-Mediated Settings

To unlock the transformative potential of LLMs in writing instruction, feedback must be reimagined not merely as correction, but as a literacy—an evolving set of interpretive, reflective, and dialogic

competencies that learners actively develop. Feedback literacy, as defined by Carless and Boud (2018: 1319), involves understanding feedback processes, recognising its role in learning, and acting on it effectively. In GenAI-mediated classrooms, this literacy must extend to *algorithmic feedback discernment*.

Accordingly, feedback literacy must be complemented by *AI literacy*: the capacity to understand how GenAI systems generate feedback, recognise their probabilistic nature and limitations, and distinguish algorithmic suggestions from pedagogically intentional guidance. AI literacy intersects with, but does not replace, feedback literacy. Without this layered competence, learners risk treating pragmatically hedged, probabilistic AI output as authoritative correction, defaulting to passive uptake rather than reflective engagement.

Developing authorship awareness, therefore, requires learners to interpret, negotiate, and selectively adopt GenAI feedback. Although LLMs can provide flexible, context-sensitive suggestions, their outputs may be generic, misaligned with rhetorical intent, or culturally inappropriate. Learners must also develop pragmatic discernment: recognising when AI reformulations alter stance, modulate politeness, intensify or soften claims, or disrupt discourse expectations relative to local classroom norms, exam genres, or target communities of practice.

Crucially, this process is mediated by the *Affective-Epistemic Feedback Habitus* learners bring to the classroom. Where learners have been socialised into authority-driven, correction-oriented feedback regimes, dispositions of deference, anxiety, and compliance may lead to uncritical acceptance of AI suggestions, undermining dialogic engagement unless explicitly addressed.

To embed LLMs into feedback processes rather than products, teachers can design tasks that require learners to compare AI suggestions with original intentions, justify acceptance or rejection of revisions, annotate changes using genre- or register-based rationales, and reflect on how feedback reshapes rhetorical stance. Such practices position GenAI feedback as a catalyst for reflection rather than an endpoint for correction, aligning with constructivist principles of learner agency.

Finally, explicit classroom discussions about AI's limits, its biases, lack of socio-pragmatic awareness, and tendency toward over-editing can demystify the tool and foster critical digital literacy. Students thus come to see AI not as an infallible authority, but as a helpful, fallible, and negotiable feedback partner. These discussions should include "LLM pragmatics": how prompts function as speech acts, how politeness and stance are conveyed in AI feedback, and how culturally variable interpretations of hedging, directness, and evaluation can shape learner uptake. This is where AI literacies intersect with feedback literacies: the goal is not merely competent use of tools, but informed judgment about when AI supports learning and when it substitutes for it.

4.3 Localising GenAI Use in Writing Feedback Across Educational Contexts

Crucially, the successful integration of LLMs into feedback practices requires culturally responsive adaptation across educational contexts. Writing classrooms—across EFL, ESL, and academic literacy settings—operate within sociocultural ecosystems where feedback is tightly intertwined with identity, authority, evaluation, and emotional exposure. Ignoring these dynamics in favour of universalised "best practices" risks reproducing the very disconnects that have historically limited the effectiveness of feedback in diverse educational systems (Vogt & Tsagari, 2014: 59). Localisation, therefore, should not be misconstrued as parochialism: it constitutes a core design principle for responsible GenAI integration, precisely because feedback is always situated, even when the technologies that deliver it are global.

Accordingly, GenAI must be localised rather than merely adopted. This entails mediating LLM use in ways that account for learners' linguistic profiles, rhetorical traditions, and socio-affective expectations. Across many educational contexts, learners bring specific patterns of L1 transfer, genre socialisation, and assessment-oriented writing habits that shape how feedback is interpreted and acted upon. AI-generated feedback that simply flags errors or rewrites text without contextual explanation risks reinforcing surface-level correction, regardless of the language or setting. Educators, attuned to both the linguistic and emotional realities of their classrooms, are uniquely positioned to provide this mediation, ensuring that AI feedback supports learning rather than compliance. Localisation must also involve pragmatic calibration: expectations surrounding directiveness, mitigation, praise, and critique vary widely across educational cultures, and LLM outputs may inadvertently signal excessive authority, unwarranted tentativeness, or ambiguity. Attending the pragmatics of LLM feedback is therefore integral to culture-sensitive feedback design.

At the same time, GenAI integration must actively resist techno-solutionism, the assumption that automation, speed, or efficiency alone can resolve pedagogical challenges. When LLMs are introduced primarily to address teacher workload, turnaround time, or error density, they **risk** digitising correctional paradigms rather than transforming them. In exam-oriented and accountability-driven systems worldwide, such adoption frames GenAI as a tool for rapid surface-level optimisation, thereby reinforcing behaviourist feedback logics under the guise of innovation. Teacher-guided adaptation is therefore essential: educators must deliberately curate how and when GenAI is used, monitor its effects on learner motivation and revision behaviour, and embed it within feedback cycles that privilege process, interpretation, and reflection over performance outcomes.

Addressing these risks requires sustained teacher professional development that extends well beyond technical training. Educators need structured support in critical AI literacy, pedagogical mediation, and affect-sensitive feedback design, enabling them to interpret GenAI outputs, anticipate learner misinterpretations, and integrate AI feedback into dialogic instructional routines rather than allowing it to function as a parallel or substitute authority. Professional development should explicitly address AI literacies, including tool evaluation, prompt pedagogy, authorship, and ethical considerations, as well as pragmatic awareness, such as how AI feedback communicates stance, politeness, and implied authority across contexts.

Reimagining feedback cultures through GenAI, therefore, demands more than functional integration; it requires a cultural, emotional, and pedagogical recalibration across educational systems. Teachers must be repositioned as mediators of meaning, learners cultivated as literate feedback agents, and AI grounded in the socio-affective realities of local classrooms. Only under these conditions can feedback evolve from correction to collaboration, from surveillance to scaffolding, and from red ink to intelligent dialogue.

5. Conclusion and Future Directions

As LLM-mediated tools increasingly permeate educational settings, their implications for feedback practices in EFL writing demand urgent theoretical and pedagogical scrutiny. This article has argued that GenAI is not a neutral instructional aid, but a socio-technical actor whose effects are shaped by the pedagogical, cultural, and institutional ecologies into which it is introduced. Across EFL contexts, GenAI can either reinforce entrenched correction-dominant paradigms or contribute to the reconfiguration of feedback cultures as dialogic, reflective, and learner-centred processes. Crucially, where feedback has historically been organised around behaviourist logics of error detection, compliance, and external evaluation, uncritical adoption of GenAI risks assimilating new technologies into existing correctional regimes rather than transforming them.

A central risk identified in this paper is techno-solutionism: the assumption that the introduction of advanced technologies will, in and of itself, resolve long-standing pedagogical problems related to feedback quality, learner engagement, or teacher workload. When GenAI is adopted primarily for its efficiency, speed, or scalability, without explicit pedagogical reframing, it risks functioning as a *digital red pen*, automating correction-as-judgment and intensifying compliance-oriented revision practices through algorithmic means. Techno-solutionism thus represents not a separate issue but a *mechanism of reinforcement*: it enables behaviourist feedback logics to persist under the appearance of innovation, particularly in assessment-driven EFL systems where surface accuracy remains the most visible indicator of improvement.

The analysis underscores that resisting techno-solutionism requires more than caution; it requires explicit cultivation of literacies. First, feedback literacy is essential if learners are to engage with GenAI feedback as an interpretive resource rather than a verdict. Without the capacity to question, negotiate, and justify revisions, learners may default to passive uptake, reproducing epistemic dependency rather than developing authorial agency. Second, AI literacies are indispensable to understanding how LLMs generate feedback, recognising probabilistic and pragmatically hedged outputs, and evaluating suggestions against rhetorical intent, task criteria, and audience expectations. These literacies are not peripheral but foundational: without them, GenAI risks becoming an invisible curriculum that silently reshapes what counts as “good writing” and “successful revision.”

Equally critical is the role of teacher professional development in mediating these processes. Teachers remain central to preventing uncritical GenAI adoption, not only by guiding tool use but by framing feedback as a dialogic, affectively safe, and culturally situated practice. Professional development must therefore extend beyond technical training to include critical AI literacy, pedagogical mediation, ethical awareness, and pragmatic sensitivity to how AI feedback communicates stance, authority, and evaluation. Without such preparation, educators may inadvertently legitimise correction-only uses of GenAI or cede epistemic authority to algorithmic outputs, allowing techno-solutionism to reconfigure, not reduce, existing feedback asymmetries.

This theoretical inquiry opens multiple pathways for future empirical research across EFL contexts. There is a pressing need for longitudinal studies examining how AI-mediated feedback shapes revision depth, authorship awareness, emotional engagement, and metacognitive development over time. Further research should also investigate pragmatic alignment in LLM feedback, specifically how learners interpret hedging, directiveness, politeness, and stance, and how these cues interact with local norms of authority and evaluation. Importantly, future studies should foreground student voice, exploring how learners experience GenAI feedback emotionally and epistemically, and how these experiences influence uptake, resistance, or overreliance.

The integration of GenAI into EFL writing education should therefore be approached not as a technological upgrade but as a pedagogical and ethical inflection point. Whether GenAI amplifies dependency or fosters autonomy, entrenches correctional authority or redistributes interpretive agency, depends less on algorithmic sophistication than on the literacies, mediation practices, and professional judgments that surround its use. Feedback ecosystems must be intentionally designed, where teachers mediate, learners interpret, and AI functions as dialogic support rather than digital overseer.

Only under these conditions can red ink be transformed into resonance and correction into collaboration. Ultimately, reshaping feedback in EFL education requires not only new tools but the unlearning of the established *Affective-Epistemic Feedback Habitus*, a pedagogical challenge that demands cultural sensitivity, teacher mediation, and dialogic reorientation. The challenge ahead is therefore not merely to

teach with AI, but to teach against the grain of inherited feedback cultures, crafting classroom ecosystems where both human and algorithmic feedback can be sources of empowerment, not surveillance. The question we must now confront is simple, yet consequential: Will we use GenAI to reinforce the red pen or to rewrite the culture of writing itself?

Limitations

This article is positioned as a theoretical and critical discussion rather than an empirical investigation or systematic review. As such, its claims are conceptual and interpretive, grounded in established scholarship, theoretical synthesis, and contextual analysis rather than primary data collection. The illustrative examples employed serve explanatory purposes, not evidentiary ones. Accordingly, the arguments advanced here should be understood as a theoretically informed agenda for future research and pedagogical reflection, intended to be empirically examined, refined, and contextualised through methodologically rigorous studies in EFL writing classrooms and comparable educational settings.

References

- Ajjawi, R., & Boud, D. (2017a). Exploring the tensions of giving and receiving feedback in higher education. *Studies in Higher Education*, 42(4), 584–597. <https://doi.org/10.1080/02602938.2015.1102863>
- Ajjawi, R., & Boud, D. (2017b). Researching feedback dialogue: An interactional analysis approach. *Assessment & Evaluation in Higher Education*, 42(2), 584–597.252–265. <https://doi.org/10.1080/02602938.2015.1102863>
- Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency* (pp. 610–623). Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>
- Bourdieu, P. (1977). *Outline of a theory of practice*. Cambridge University Press.
- Brown, P., & Levinson, S. C. (1987). *Politeness: Some universals in language usage*. Cambridge University Press.
- Carless, D., & Boud, D. (2018). The development of student feedback literacy: Enablinguptake of feedback. *Assessment & Evaluation in Higher Education*, 43(8), 1315-1325. <https://doi.org/10.1080/02602938.2018.1463354>
- Carless, D., & Winstone, N. (2020). Teacher feedback, literacy, and student engagement with feedback: Reviewing the evidence. *Educational Research Review*, 31, 100326. <https://doi.org/10.1016/j.edurev.2020.100326>
- Cavaleri, M., Jenkins, K., & Somasundaram, J. (2024). Harnessing generative AI for academic writing: Affordances, anxieties, and pedagogical strategies. *Computers and Composition*, 69, 102801. <https://doi.org/10.58723/jaiela.v1i1.56>
- Danescu-Niculescu-Mizil, C., Sudhof, M., Jurafsky, D., Leskovec, J., & Potts, C. (2013). A computational approach to politeness with application to social factors. In *Proceedings of the 51st Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 250–259). Association for Computational Linguistics.
- Dewaele, J.-M. & Li, C. (2020). Emotions in Second Language Acquisition: A Critical Review and Research Agenda. *Foreign Language World*, (1): 34-49.
- Floridi, L., Cowls, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People—An ethical framework for a good AI society: Opportunities, risks, principles, and recommendations. *Minds and Machines*, 28(4), 689–707. <https://doi.org/10.1007/s11023-018-9482-5>
- Garcia, A., & Lee, M. (2023). Learning with machines: Student attitudes and practices in AI-supported writing. *Language Learning & Technology*, 27(2), 85–97.

- Horwitz, E. K., Horwitz, M. B., & Cope, J. (1986). Foreign language classroom anxiety. *The Modern Language Journal*, 70(2), 125–132.
- Hovy, D., & Spruit, S. L. (2016). The social impact of natural language processing. In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)* (pp. 591–598). Association for Computational Linguistics. <https://doi.org/10.18653/v1/P16-2096>
- Lantolf, J. P., & Thorne, S. L. (2006). *Sociocultural theory and the genesis of second language development*. Oxford University Press.
- Liu, X., & Wang, H. (2023). Generative AI and writing development: From language correction to authorship support. *Language Learning & Technology*, 27(1), 103–121.
- Lyriqkou, C. (2021). *The role of informal second language learning in the spoken productions of EFL learners in Greece* (Doctoral dissertation). The Open University, UK.
- Nicol, D. (2021). The power of internal feedback: Exploiting natural comparisons in assessment design. *Assessment & Evaluation in Higher Education*, 46(3), 450–462. <https://www.tandfonline.com/doi/full/10.1080/02602938.2020.1823314>
- Ryan, T., & Henderson, M. (2018). Feeling feedback: Students' emotional responses to educator feedback. *Assessment & Evaluation in Higher Education*, 43(6), 880–892. <https://www.tandfonline.com/doi/full/10.1080/02602938.2017.1416456>
- Semke, H. D. (1984). Effects of the red pen. *Foreign Language Annals*, 17(3), 195–198.
- Ushioda, E. (2011). Language learning motivation, self, and identity: Current theoretical perspectives. *Computer Assisted Language Learning*, 24(3), 199–210. <https://doi.org/10.1080/09588221.2010.538701>
- Vlanti, S. (2012). Assessment practices in the English language classroom of Greek Junior High School. *Research Papers in Language Teaching and Learning*, 3(1), 92–122
- Vogt, K., & Tsagari, D. (2014). Assessment literacy of foreign language teachers: Findings of a European study. *Language Assessment Quarterly*, 11(4), 374–402. <https://doi.org/10.1080/15434303.2014.960046>
- Vygotsky, L. S. (1978). *Mind in society: The development of higher psychological processes*. Harvard University Press.
- Winstone, N. E., & Carless, D. (2020). *Designing effective feedback processes in higher education: A learning-focused approach*. Routledge.

Ioanna Nifli-Sakali (iniflis@enl.auth.gr) is a Greek-Canadian English language educator, translator, and Ph.D. candidate at the School of English, Aristotle University of Thessaloniki. She holds an MA in TESOL and an MA in Applied Translation Studies from the University of Leeds. Her professional experience includes high-stakes institutional work with international organisations such as the United Nations and the European Parliament, supporting complex policy-driven communication across multilingual environments. Her doctoral research in Computational Linguistics examines the pedagogical implications of Large Language Models (LLMs) and Automated Essay Scoring (AES) systems in Greek EFL writing contexts, with broader interests in AI-mediated language learning, multilingual communication, and responsible educational innovation.