

AI-mediated feedback in English learning: Moroccan graduate students' accounts of motivation, autonomy, and trust conditions

Nouh Alaoui Mhamdi

Sidi Mohamed Ben Abdellah University, Faculty of Letters and Human Sciences, Dhar El Mahraz, English Department, Fez, Morocco

ARTICLE INFO

Keywords:

AI-mediated feedback
Technology acceptance model
Feedback uptake
Assessment for learning
Moroccan EFL learners

ABSTRACT

This article examines how Moroccan graduate students describe using generative-AI tools (e.g., ChatGPT) as sources of feedback for English learning, with attention to motivation and autonomy as enacted through feedback uptake and trust calibration. Sixty-two graduate students at Sidi Mohamed Ben Abdellah University (Fez) completed an online qualitative survey (Oct 2024-Feb 2025) comprising seven open-ended questions; prompts were delivered in French, and responses were provided in French, Arabic, or English, then translated into English under a meaning-preservation rule and analysed through reflexive thematic coding in MAXQDA 24. Using the Technology Acceptance Model as a heuristic frame, perceived ease of use is linked to rapid access and low-friction interaction, while perceived usefulness is described as condition-dependent across task stakes, access stability, and pragmatic-cultural fit. Participants frame "improvement" through task-situated performance dimensions (accuracy, fluency, controlled complexity, genre/register appropriacy) and through a feedback-use cycle in which tasks are specified, outputs inspected for fit, weaknesses diagnosed, prompts refined, and suggestions then adopted, edited, verified, or rejected. The analysis supports an assessment-for-learning interpretation in which educational value is associated with mediation routines and calibrated trust, rather than automated judgement.

1. Introduction

Recent advances in artificial intelligence (AI) are associated with shifting practices in higher education language learning, particularly through tools that generate text, explanations, and feedback on demand. In language education contexts, generative systems are commonly taken up as writing and revision supports (e.g., drafting, paraphrasing, summarisation) and as interactional resources that can extend opportunities for practice [1,2]. At the same time, the instructional value of AI-mediated feedback is frequently treated as conditional: rapid output and surface-level correctness can align with efficiency, yet they can also align with risks of procedural substitution when learners accept outputs without evaluation [3,4]. These tensions are likely to be especially salient in multilingual settings where English learning occurs alongside other instructional languages and where access constraints can shape who can iterate, verify, and sustain feedback use over time [5,6]. In Morocco, these conditions are relevant for understanding how AI may relate to learning opportunities, trust, and equity in graduate-level English use.

Empirical work on AI in language learning often links adoption to

immediate feedback and flexible practice, including non-judgmental interaction that can support willingness to engage in oral rehearsal [1] and drafting support that can expand options for phrasing and lexical choice [7,2]. However, contemporary language education research commonly frames "learning" as multi-dimensional and process-oriented: it relates not only to outputs that appear correct, but also to how learners use language (e.g., accuracy, fluency, controlled complexity), how they take up feedback, and how they regulate effort across tasks. In this scholarship, autonomy is often treated as conditional on evaluation practices: instant feedback can align with self-paced study when learners inspect, revise, and verify, while routine acceptance of outputs can align with lower critical evaluation and weaker self-regulation [4,8]. These patterns position feedback uptake and trust calibration as central to understanding what AI support means in practice, especially in academically consequential writing contexts.

Socioeconomic factors further complicate AI integration in language learning. Learners from lower-income backgrounds are often reported to face difficulties accessing advanced AI features because premium services remain financially prohibitive [9,10]. This situation has been associated with disparities in the quality of feedback and support that

E-mail address: alaouimhamdi.nouh@usmba.ac.ma.

<https://doi.org/10.1016/j.caeo.2026.100356>

Received 18 May 2025; Received in revised form 1 April 2026; Accepted 10 April 2026

Available online 12 April 2026

2666-5573/© 2026 The Author(s). Published by Elsevier Ltd. This is an open access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

students receive [6]. These patterns suggest that digital inequality may limit AI's potential to support more equal educational opportunities and point to the need for closer examination of how AI can either mitigate or reproduce existing inequities [5].

This study examines how Moroccan graduate students describe using AI tools for English learning, with attention to motivation and autonomy as enacted through feedback uptake and trust calibration. Prior work has documented adoption, motivation, and autonomy in AI-supported language learning, yet a qualitative gap remains concerning how learners in multilingual, resource-constrained settings judge AI feedback in practice, when they trust it, and where they set boundaries. The Technology Acceptance Model (TAM) [11] is used heuristically, with perceived usefulness (PU) and perceived ease of use (PEOU) treated as situated evaluations shaped by task demands, access constraints, and pragmatic-cultural fit.

Accordingly, the study addresses the following questions:

RQ1. What English-learning targets do participants report when using AI, in terms of skill domains (e.g., writing/genre, grammar/accuracy, vocabulary/phraseology, reformulation/translation, speaking/interaction) and performance emphases (accuracy, fluency, controlled complexity, genre/register appropriacy)?

RQ2. What feedback functions and uptake practices are reported (e.g., correction, modelling, explanation, practice generation; iterative prompting, verification, restriction), and how do these practices constitute a feedback-use cycle?

RQ3. How do access constraints and pragmatic-cultural fit relate to reported usefulness and trust calibration, and what practice-based uptake profiles are visible in participants' accounts?

RQ4. How do participants position AI relative to teachers in English learning, and what safeguards and classroom routines are implied by their accounts?

Accordingly, the study contributes qualitative evidence from a Moroccan graduate cohort based on open-ended survey responses (N = 62), clarifying how PU and PEOU relate not only to stated usefulness, but also to feedback use, verification, restriction, and selective trust. In this respect, the article extends earlier work by linking adoption-related judgements to concrete learner practices in a multilingual Moroccan context and by specifying classroom-relevant trust conditions rather than adding another general account of AI acceptance. It also offers a standards-informed discussion of language ability domains (e.g., accuracy, coherence, pragmatic appropriacy) as interpretive anchors for pedagogical implications, without treating the dataset as a validation study of automated scoring.

2. Literature review

This review is organised into four strands: adoption and motivation, autonomy and learning quality, equity and access, and TAM in language learning. Across these strands, usability gains are widely reported, but learning quality, trust, equity, and contextual fit remain conditional. What remains under-specified is how learners in multilingual, resource-constrained settings judge AI-mediated feedback in practice, when they trust or restrict it, and how those judgements relate to feedback uptake. This review therefore synthesises these strands to define the gap addressed here: a qualitative account of Moroccan graduate learners' feedback uptake, trust calibration, and boundary-setting, interpreted through TAM as a situated rather than universal frame.

2.1. Adoption and motivation

Studies report that AI tools can raise motivation through rapid, low-stakes feedback and flexible practice. Chatbots support speaking without fear of judgement, and writing assistants help learners try new forms and extend vocabulary [9,7,1,2]. Recent work in EFL oral reading reports that motivation correlates strongly with reading engagement and also, to a lesser degree, with interaction with AI tools, with

differences between undergraduate and master's students [12]. At the same time, several papers warn that outsourcing planning or relying on surface edits risks shortcut habits that may weaken effort on complex tasks [3]. Taken together, prior findings suggest that motivation tends to rise when AI is channelled toward active production and feedback, and appears to stall when tools are used primarily as a shortcut. In this study, "adoption" is read as motivated uptake that sustains self-directed work rather than one-off reliance on corrections.

2.2. Autonomy and learning quality

Findings on autonomy are mixed. Instant explanations and always-on feedback can support self-paced study and confidence [8]. Other studies note risks to self-regulation and critical judgement when outputs are accepted uncritically [2,4,13]. Outcome-focused research proposes a mechanism: in Pallant et al.'s (2025) study, when students used generative AI to construct and extend knowledge, higher-level learning appeared to improve, whereas more procedural use seemed to be associated with lower-level outcomes. In writing, AI feedback tends to raise sentence-level accuracy and variety, while hybrid feedback that combines teacher and AI produces stronger gains in organisation and argument and is preferred by learners [6,14]. The review therefore treats autonomy as both supported and threatened by AI, and examines when constructive use is reported versus procedural substitution that may be associated with weaker self-regulation. In this regard, autonomy is likely to be supported when AI scaffolds planning, drafting, and revision, and appears threatened when it substitutes for those steps. A practical corollary often drawn in this literature is that instructors may wish to retain human input where higher-order goals matter and use AI primarily to stage practice and iteration.

2.3. Equity and access

Access conditions appear to shape who benefits. Paywalls and unstable connectivity can widen gaps in support and feedback quality [10,5]. Moroccan evidence on flipped learning reports higher engagement and self-direction but also uneven digital skills and infrastructure, with calls for training and institutional support [15]. These patterns suggest that affordability, bandwidth, device access, and staff development are likely to condition the value students attach to AI support in this context. Affordability, bandwidth/connectivity, and device access are treated as moderators of PU, and the analysis examines whether free tiers are described as sufficient for routine study tasks while premium features remain gated.

2.4. Cultural and contextual accuracy

Language tasks that require cultural nuance appear to expose limits in some AI outputs. Reports note tone, register, and rhetoric issues when systems follow Anglophone norms or miss discourse cues such as irony, which can distort intended meaning in non-Western settings [16,7,5,2]. For multilingual cohorts that work in French and Arabic as well as English, these limits can dampen PU on advanced writing tasks. Culturally tuned prompts, local exemplars, and teacher mediation can mitigate drift, although they are unlikely to eliminate it entirely. Given routine use of French and Arabic in this cohort, cultural and linguistic fit is treated as a condition on usefulness for advanced writing rather than an afterthought.

2.5. Technology acceptance model in language learning

TAM links adoption to PU and PEOU [11]. In language and education studies, PEOU is often reported to support early uptake, while perceived usefulness tends to be associated with continued use [17]. Recent qualitative work shows how stakeholders judge usefulness and ease in situated ways. Faculty describe benefits and limits of virtual coursework

that shape acceptance beyond tool novelty [9,18]. In assessment, students and teachers value efficiency and personalisation yet worry about explainability, trust, and the teacher-student relationship [19]. Sector overviews argue that policy, ethics, and integrity guardrails condition acceptance as much as interface factors [20]. For the present study, TAM is used heuristically rather than predictively: PU/PEOU frame interpretation, while cost, cultural fit, and assessment policy are treated as moderators of intention to use. This stance matches a reflexive qualitative design.

Few studies centre Moroccan graduate learners with qualitative evidence tied explicitly to TAM and to cultural, linguistic conditions. This review therefore suggests that a gap remains at the intersection of adoption, learning quality, equity, and cultural fit. The study addresses it by analysing PU and PEOU alongside reported benefits, risks, and access limits in a multilingual cohort.

Table 1 Summarises strands, representative evidence, and how each strand informs the present analysis.

Prior work charts motivation gains, autonomy risks, equity barriers, and acceptance conditions. What remains under-specified is a qualitative account from Moroccan graduate learners that treats TAM as situated rather than universal. This study takes up that task.

3. Methodology

This study adopts a qualitative research design to examine Moroccan graduate students' experiences with AI tools in English language learning. The analysis draws on open-ended survey responses gathered from October 2024 to February 2025 and analysed thematically in MAXQDA 24. In keeping with this design, the study focuses on interpretive analysis of participant accounts rather than statistical aggregation. The seven survey prompts are listed verbatim in Appendix A. All descriptive counts use the unit "participant"; an item is treated as present if a participant advances it at least once across Q1 to Q7 (N = 62). Multiple mentions by the same participant do not increase the count, so these counts function only as descriptive indicators of salience rather than as statistical estimates.

3.1. Research design

A qualitative design underpins this study because the research questions concern how Moroccan graduate students describe using AI-mediated feedback, what meanings they attach to it, and under what conditions they trust, verify, restrict, or reject it. The design is therefore interpretive rather than hypothesis-testing and is suited to capturing situated accounts of motivation, autonomy, and feedback uptake in a multilingual context. Reflexive thematic analysis was adopted to examine patterns in these accounts, while the TAM served as a heuristic organising frame rather than a predictive model or a full theory of language learning. Within this frame, PU and PEOU were treated as

context-dependent judgements shaped by task demands, access conditions, and pragmatic-cultural fit. Rigour was supported through an audit trail, reflexive memoing, peer debrief, member feedback, and negative-case checks rather than inter-rater coefficients, in line with the study's reflexive qualitative design.

3.2. Sampling procedure

Participants were 62 Moroccan graduate students (2024/2025 academic year) from Sidi Mohamed Ben Abdellah University in Fez, recruited through convenience sampling via mass email to complete an online qualitative survey. This approach was intended to capture varied accounts of AI-supported English learning within a single cohort. Eligibility required prior use of AI tools for English learning, and no incentives were offered. Data were gathered from October 2024 to February 2025.

3.3. Participants

A brief profile of the cohort is provided in Table 2. No demographic fields (e.g., age, gender, field of study, self-reported socioeconomic status, or proficiency scores) were collected, in order to preserve anonymity in an online format. As a consequence, the analysis does not undertake subgroup comparisons and treats all findings as situated within this single cohort.

3.4. Inclusion and exclusion criteria

Participation is open to graduate students actively enrolled at Sidi Mohamed Ben Abdellah University who have engaged with AI tools in the context of English language learning. The study excludes those who have not used such tools, so that the data more directly reflect experiences relevant to the research focus. No demographic fields were collected to protect anonymity. Screening removed empty, duplicate, or

Table 2
Participants (N = 62).

Characteristic	Description
Institution	Sidi Mohamed Ben Abdellah University, Fez - Morocco
Faculty	Faculty of Letters and Human Sciences, Dhar Mahraz
Level	Graduate students (Master's / doctoral programmes)
Academic year	2024–2025
Inclusion criterion	Prior use of AI tools for English language learning
Recruitment	Convenience sample via mass email invitation
Data collection window	October 2024 to February 2025
Response languages	English, French, Arabic (translated to English for analysis)
Demographic variables	Not collected (no age, gender, major, socioeconomic status, or proficiency)

Table 1
Evidence map for AI in English language learning in higher education.

Strand	Core claim	Representative evidence	Implication for this study
Adoption & motivation	AI tools can raise motivation through rapid, low-stakes feedback and flexible practice; effects vary by level and task.	[1]; Alfaleh, 2025	Expect motivation gains where tools support active production, check differences across graduate tasks.
Autonomy & learning quality	Benefits for self-paced study, but risks to self-regulation if outputs are accepted uncritically; learning depth depends on how AI is used.	Zou, 2025; Shi, 2023; Hong, 2024; Pallant, 2025	Emphasise constructive, mastery-oriented use, caution against procedural substitution.
Writing development & feedback	AI improves sentence-level accuracy; hybrid human+AI feedback outperforms AI-only on organisation and argument.	[6]; Zhang, 2025	Recommend hybrid feedback where higher-order goals matter.
Equity & access	Paywalls and connectivity widen gaps; local skills and infrastructure mediate benefits.	[5]; Mjihad, 2025	Treat affordability, bandwidth, device access, and training as conditions on usefulness.
Cultural/contextual fit	Anglophone norms can misread tone and rhetoric in non-Western settings; risk of cultural mismatch.	[16]	Use culturally tuned prompts, local exemplars, and teacher mediation, interpret usefulness as conditional.
Acceptance (TAM)	PEOU supports early uptake; PU sustains use; trust, explainability, and policy shape acceptance.	[18,11,17,20,19]	Interpret findings through PU/PEOU while naming cost and culture as moderators.

off-topic entries (e.g., repeated "none"), using timestamp and identical-string checks across items. Because no demographic fields were collected, subgroup contrasts are not attempted, and the analysis is presented at the cohort level.

3.5. Ethical considerations

The survey was anonymous and collected no identifying information. Participants provided informed consent via an online information sheet and tick-box prior to participation. No demographic fields were collected to reduce re-identification risk, and all cases were assigned random alphanumeric codes (e.g., P01). Data were stored in password-protected project files with local backup, and only de-identified excerpts were used in reporting; non-English excerpts were translated into English. The study was conducted under local institutional conditions in which formal ethics review structures are limited; confidentiality and minimal-risk procedures were therefore prioritised through anonymity, restricted data handling, and careful excerpt selection.

3.6. Data collection procedures

Data were collected through an anonymous asynchronous online qualitative survey comprising seven open-ended questions. This mode was chosen to permit detailed narrative responses while reducing the perceived exposure associated with discussing generative-AI use in academic settings, where some students may have been hesitant to participate in interview-based formats. The survey therefore prioritised privacy, flexibility, and self-paced completion over live interaction; no probing or rapport-building occurred, but response depth was supported through open-ended prompts and the option to answer in French, Arabic, or English. Prompts were delivered in French, and non-English responses were translated into English before coding under a meaning-preservation rule (literal where possible; idioms glossed in brackets), followed by a short translation check during cleaning. The full prompt list appears in Appendix A. Project files were stored in MAXQDA 24 in a versioned folder on a university Google Drive account, with a password-protected local backup.

3.7. Data analysis

Data were analysed through reflexive thematic analysis in MAXQDA 24 to produce an interpretive account of AI-mediated feedback use in English learning tasks. Coding proceeded iteratively at meaning-unit level and was supported by a living codebook with definitions, inclusion/exclusion criteria, and anchor excerpts. TAM functioned as a sensitising frame for adoption-related statements, while reporting was organised around five dimensions visible in the corpus: skill targets, performance emphases, feedback functions, uptake practices, and trust risks/failure modes. Descriptive counts follow a participant-level mention rule and are reported only as indicators of salience.

3.8. Analytic pipeline and rigour

1. Import responses into MAXQDA 24 and create a case for each participant; attach question identifiers (Q1-Q7) to each response segment.
2. Open a reflexive memo documenting analytic stance and the role of TAM as a sensitising frame for interpreting PU and PEOU.
3. Draft Codebook v0 from the research aims and an initial corpus read; define inclusion/exclusion rules and anchor excerpts for each code.
4. Pilot-code approximately ten percent of the corpus line by line; log decision points (code boundaries, overlaps, and ambiguous segments) in a decisions memo.

5. Revise to Codebook v1; merge duplicates, split over-broad codes, and clarify rules for distinguishing uptake practices from outcome claims.
6. Code the full corpus line by line with constant comparison; add in-vivo codes where participant phrasing is analytically salient.
7. Organise codes into five reporting dimensions aligned with the Findings: skill targets; performance emphases (accuracy/fluency/controlled complexity/genre appropriacy); feedback functions; uptake practices; trust risks/failure modes.
8. Generate descriptive participant-level summaries for each dimension (presence/absence per participant) and export these to support the construction of tables (participant-level mention counts).
9. Stress-test boundaries between closely related categories (e.g., verification vs restriction; meaning drift vs contextual misfit; efficiency claims vs dependency concerns) using negative-case checks and memoed rationale.
10. Select representative excerpts and counter-excerpts for each analytic claim within each dimension, ensuring coverage across prompts (Q1-Q7) rather than single-item dependence.
11. Produce integrative analytic summaries that connect dimensions (e.g., how access constraints reshape iteration; how pragmatic-cultural fit relates to verification/avoidance; how uptake profiles align with different scaffolding needs), then check alignment with the research questions and the TAM frame.

3.9. Credibility and trustworthiness

Credibility and trustworthiness were supported through an audit trail, reflexive memoing, peer debrief, member feedback, and negative-case checks. Credibility related to iterative coding, the use of representative and counter-excerpts, and follow-up member feedback from invited participants; transferability was supported through explicit description of the institutional setting, cohort, data-collection mode, and multilingual response context. Dependability and confirmability were addressed through a documented coding trail, versioned analytic decisions, and reflexive notes on possible interpretive bias arising from contextual familiarity with the educational setting. A peer debrief helped refine category boundaries, especially between dependency risk and procedural substitution, while negative and divergent cases were retained to qualify central claims. Single-coder analysis is acknowledged; within a reflexive thematic analysis design, trustworthiness rests on transparency, documentation, and critical self-audit rather than inter-rater coefficients.

4. Findings

This section reports patterns in Moroccan graduate students' open-ended accounts of using generative-AI tools for English learning (N = 62). The reporting logic is assessment-for-learning: what learners treat as improvement (e.g., accuracy, fluency, controlled complexity, genre appropriacy), how feedback is requested and taken up, and how trust and task selection are conditioned by access constraints and contextual fit. Mention counts function only as descriptive signposts of salience within this corpus; interpretive weight is carried by the qualitative evidence, including contrasts and boundary cases.

4.1. What "improvement" means in practice: skill targets and performance dimensions

Across responses, "learning" is rarely framed as a score. Instead, participants describe improvement through performance dimensions that are repeatedly invoked in their task accounts: accuracy (grammar, mechanics, error correction), fluency (naturalness and flow), and controlled complexity (organisation, register, lexical range suited to academic and professional genres). These dimensions are not presented

abstractly; they are linked to concrete practices such as revising drafts, reformulating meaning while attempting to preserve intent, rehearsing interaction through repeated exchanges, and producing institutionally acceptable writing for coursework and professional communication. The corpus therefore depicts "improvement" as task-situated performance refinement and feedback uptake, rather than as a single metric of achievement.

Several participants frame AI as a tool for shaping academic and professional discourse. One participant describes a combined effect on structure, lexical choice, and correctness in institutional genres:

"Using tools such as ChatGPT has changed how academic and professional English texts are produced. It helps structure ideas, enrich vocabulary, correct grammatical errors, and write sentences that sound more natural. In academic work, it helps reformulate arguments clearly and concisely, and in professional settings it helps draft emails or reports with an appropriate tone." (P21)

Another participant offers a task-based account that links genre work to efficiency and to a perceived risk of reduced reflection:

"It helps me write professional emails and reformulate my messages in a professional way. The main advantage is that it saves time, but the drawback is that it can make people forget creativity and reflection. It seems important to learn how to work properly with AI tools in order to obtain accurate answers." (P25)

Alongside these benefits, a minority frames the same "smoothness" as a stylistic risk. One participant links fluency gains to perceived uniformity and reduced individuality:

"It helps with clarity and grammatical correctness, but my texts can become too uniform. It tends to structure sentences in recurring ways, which can make the writing feel artificial. For that reason, I try to keep a personal touch rather than accept every reformulation as final." (P61) [Table 3](#).

Translation and reformulation are often described as "intelligent" when they preserve intended meaning and adapt phrasing to disciplinary constraints. One participant frames this as meaning-oriented rewriting rather than literal translation:

"It reformulates my original text in a way that is simple, understandable, and focused. It reformulates according to the topic; it does not translate mechanically like Google Translate. Sometimes it even corrects the meaning, since it draws on a wide range of information." (P08)

A second participant describes domain-specific translation as an accuracy-and-register task, while also naming a recurrent failure: meaning drift.

"It helps translate texts in a way that fits the field, using the technical terms of that field. That helps me understand the text better. But sometimes the reformulations can deviate from the original meaning, so I remain cautious when I use them in academic contexts." (P05)

Speaking and interaction practice is less frequent than writing and revision, yet a visible subset treats AI interaction as a low-pressure environment for repeated oral rehearsal. One participant links this

practice to broader learning support (plans, lessons, references) while also describing a boundary against overreliance:

"I often speak to it in English to train myself to speak. It also gives me a programme for learning English, lessons, and references in English, which saves time and effort and increases my enthusiasm for learning. At the same time, it should not be trusted completely, since relying on it too much can make a person less active." (P04)

Another participant frames voice interaction as a modality shift, enabling practice beyond writing:

"When questions are asked in English, the responses also come in English, and the level of English in the prompt matters less than expected. With the voice conversation feature, it becomes more appropriate for learning English, since it allows the user to speak, not only write. It also helps turn scattered ideas into a correct paragraph when writing is required." (P28)

Taken together, these accounts position AI use as closely linked to form-focused revision (accuracy), rhetorical polishing (genre/register), and meaning management (reformulation), with a smaller but meaningful presence of oral rehearsal.

4.2. Feedback functions and uptake: how learners work with AI responses

Participants describe feedback in ways that foreground how they work with it, rather than treating it as a final judgement. Four feedback functions recur: (a) correction of form, (b) modelling and reformulation, (c) explanation, and (d) practice generation. Uptake is narrated through strategies such as iterative prompting, verification, and restriction to low-stakes uses [Table 4](#).

A representative account links multiple feedback functions to concrete learning actions, dialogue rehearsal, correction, and clarification, suggesting an iterative, self-regulatory use of feedback:

"These tools allow English improvement by simulating interactive conversations, providing simple grammar explanations, and generating tailored exercises. For example, it is possible to practise a job-interview dialogue, ask for corrections on a text, or clarify verb tense use. When the responses are generic or imprecise, I rephrase the question and specify expectations more clearly." (P21)

Verification emerges as a safeguard for credibility and appropriacy, particularly when outputs are used for study-related tasks. One participant, who otherwise used AI mainly for research support, frames inaccuracies as a recurring condition that requires checking:

"AI tools offer quick access to information and can support personalised practice, but they can be inaccurate and may encourage over-reliance. Even when used mainly for research rather than basic language learning, a general challenge is occasional inaccuracies. That can be addressed by cross-checking with reliable sources, especially when the information will be used for academic purposes." (P53)

A smaller set connects "successful use" to technical command of prompting, while still acknowledging error risk, yet the broader pattern across responses is better characterised as a feedback-use cycle that is repeated over time rather than as a single interaction that ends at output

Table 3
Reported AI-supported learning targets by skill domain (participant-level mentions, N = 62).

Skill domain (explicitly referenced)	Illustrative learning target(s)	n (%)
Writing and genre (emails, reports, academic texts)	organisation, register, coherence	30 (48.4)
Grammar, mechanics, and accuracy	error correction, tense choice, punctuation	30 (48.4)
Translation and reformulation	meaning-preserving paraphrase, summarising	28 (45.2)
Vocabulary and phraseology	lexical choice, synonymy, collocations	26 (41.9)
Speaking and interaction practice	dialogue rehearsal, voice/audio use	18 (29.0)
Input support (resources)	podcasts/videos, suggested materials	8 (12.9)

Note. Percentages indicate participant-level salience only and do not support statistical inference.

Table 4
Feedback functions and uptake stances (participant-level mentions, N = 62).

Category	Description in participant accounts	n (%)
Feedback functions		
Correction of form	grammar/spelling/punctuation correction	31 (50.0)
Models/exemplars	templates, reformulations, genre scaffolds	23 (37.1)
Explanation	rules, clarifications, metalinguistic support	15 (24.2)
Practice generation	exercises, simulated dialogue	9 (14.5)
Uptake stances		
Iterative prompting	stepwise querying; refining constraints	6 (9.7)
Verification	cross-checking with sources/teachers	9 (14.5)
Restriction/avoidance	limiting sensitive/high-stakes tasks	19 (30.6)
Dependency concern	reduced effort; routinised reliance	20 (32.3)

Note. Percentages indicate participant-level salience only and do not support statistical inference.

delivery. Participants commonly describe a sequence in which a task is specified (often with genre constraints), an output is inspected for fit (tone, coherence, lexical appropriacy), weaknesses are diagnosed (genericness, meaning drift, factual uncertainty), and the prompt is revised to narrow scope or add constraints; the revised output is then either adopted, edited manually, or subjected to verification against sources or teacher judgement when stakes are higher. In these accounts, "learning" relates to repeated movement through this cycle, prompt refinement, output evaluation, revision, and selective trust, rather than to a one-time correction event. This cyclical description provides a corpus-grounded way to represent a "learning journey" as cumulative practice in self-regulation and feedback uptake, even within a cross-sectional design.

4.3. Learner differentiation: practice-based profiles and task selection

The dataset does not include demographic variables that would support subgroup comparisons (e.g., prior proficiency, field, socioeconomic status). Variation can still be documented through practice-based profiles: what participants ask for, how they judge responses, and where they set boundaries. Four recurring profiles are visible, with overlap across individuals.

Genre-polish and efficiency users. These learners prioritise professional/academic acceptability and time economy, frequently linked to writing tasks and surface revision. A participant frames the benefits as constant availability and rapid correction, while also pointing to misunderstanding as a boundary:

"It allows me to practise English at any time, without needing a partner, and it helps organise ideas. It corrects grammar and vocabulary mistakes quickly and gives ideas and corrections fast, which saves time. But sometimes it gives incorrect answers or does not understand my questions well." (P46)

Practice-oriented oral users. These learners treat AI as a conversational partner for rehearsal and repeated attempts, especially when access to live practice is limited. Their accounts align with a process emphasis (repeated interaction) rather than a single output.

"Sometimes I chat by text, and other times I use an audio call. It helps me practise speaking and keep the conversation going when I do not have someone available to practise with." (P47)

Verification-oriented users. These accounts frame trust as conditional and treat output as provisional, especially for higher-stakes academic uses:

"I cross-check the answers with other sources, or I ask a teacher to validate them. It can support learning, but it should not be the only reference for academic work." (P21)

Boundary-setters. These learners report restricting AI use to protect authorship, effort, or credibility, and they explicitly link boundaries to graduate-level expectations:

"These tools are useful, but one should not become dependent on them, especially in academic and research contexts. It helps partially, but caution is needed when the meaning must be precise." (P06)

Counterevidence is also present. A small minority describes the tool as uniformly beneficial in the domains they use, with limited perceived risk, while still rejecting teacher replacement:

"I use it mainly in the legal field: it helps me translate legal texts into English and understand the terms used. It helped me overcome difficulties in writing, although communication and pronunciation still require time. I do not see disadvantages in my experience, but a teacher's place remains irreplaceable." (P33)

These contrasts suggest that reported risks (dependency, misfit, unreliability) are not uniformly experienced; they relate to task selection, reliance patterns, and how learners evaluate and verify outputs.

4.4. Where the system breaks down: constraints, failure modes, and learner safeguards

Participants report constraints and failure modes that shape trust and

task selection. Three clusters recur: (a) access constraints (paywalls and limits), (b) contextual misfit, and (c) unreliable or "robotic" output. Participants also describe responses to these breakdowns: restriction, verification, and iterative specification Table 5.

Access constraints are often narrated as a condition on what AI can be used for. One participant links "high-level learning" to paid access, while also acknowledging that errors remain possible:

"The outputs are usually targeted, understandable, and fluent, and the translation is more 'intelligent' than a mechanical tool. It can feel like a question/answer course, similar to what happens in a classroom. But it can still contain mistakes, since it relies on human-produced information that is not always validated. If a truly high-level kind of learning is expected, paid access becomes necessary." (P08)

Contextual misfit is most clearly articulated in sensitive domains, where perceived risk is higher. One participant ties unreliability to topics where factual precision and contextual judgement are consequential:

"It can help in many ways, listening, reading, speaking, and it can suggest sources for English learning. But it can also provide false information, especially when the topic is important: religious, historical, political, or statistical. In those cases, I cannot treat the output as reliable without checking." (P47)

Output quality concerns also include stylistic artefacts and context-insensitive replies. One participant links benefits (self-training, correction) to a recurring limitation: distorted responses that ignore context, producing "robotic" text:

"I ask the AI to correct my phrasing, expand my ideas, and correct syntactic errors, and it supports self-training. But sometimes the results are distorted or do not take the overall context of the question into account. In those cases, the generated text can feel robotic, so I do not treat it as a final version." (P60)

A recurrent tension concerns dependency and effort. Participants do not offer objective measures of learning decline, but they repeatedly frame routinised reliance as linked to reduced initiative and weaker self-directed searching habits:

"The main danger is getting used to easy access to information without effort, which can weaken the ability to search. When reliance becomes automatic, the learner can become less active and less engaged in building knowledge independently." (P20)

These accounts identify a perceived trade-off: convenience and immediate support relate to time economy and confidence, while routinised reliance relates to lower effort and weaker self-regulation.

4.5. Human teaching, AI feedback, and instructional value: substitution claims, complementarity, and classroom relevance

Participants position AI relative to teaching in three ways: rejection of replacement, conditional substitution for limited tasks, and endorsement of replacement. The dominant orientation is complementarity. Rationales clarify what learners treat as distinctively human in language education: relational support, adaptive pedagogy, and context-sensitive judgement in meaning-making Table 6.

A strong complementarity account distinguishes teacher functions

Table 5
Reported constraints and failure modes, with typical learner responses (participant-level mentions, N = 62).

Constraint / failure mode	Typical response described	n (%)
Paywalls, costs, and access limits	reliance on free tiers; narrowing to short tasks; manual drafting	23 (37.1)
Contextual/cultural misfit	avoidance of sensitive topics; teacher validation; reframing prompts	13 (21.0)
Unreliable outputs	cross-checking; limiting trust in high-stakes work	5 (8.1)

Note. Percentages indicate participant-level salience only and do not support statistical inference.

Table 6

Positions on teacher replacement (participant-level, N = 62).

Position	Typical rationale in participant accounts	n (%)
No replacement	relational support; adaptive guidance; human judgement	31 (50.0)
Conditional replacement	possible for limited functions/tasks; not fully sufficient	18 (29.0)
Replacement endorsed	availability; non-judgemental interaction; efficiency	8 (12.9)
Missing/unclear	not codable	5 (8.1)

Note. Percentages indicate participant-level salience only and do not support statistical inference.

(human interaction, cultural judgement, adaptive guidance) from AI functions (availability, immediate feedback, repeated practice without embarrassment):

"ChatGPT cannot replace an English teacher, but it can be an excellent complement. A teacher provides human interaction, emotional understanding, and cultural expertise, and adapts teaching to individual needs with motivating guidance. ChatGPT can support learning through constant availability, immediate explanations, and opportunities to practise without feeling judged." (P21)

Relational and interactional dimensions recur even in brief statements. One participant frames learning as socially constituted and treats that relational dimension as non-substitutable:

"A teacher cannot be replaced. Human relationships support language learning through reciprocal interaction, and that interaction is part of the learning itself." (P45)

Conditional substitution tends to appear when "teaching" is reduced to limited functions (correction, phrasing, templates). A participant who uses AI for professional phrasing and recognises an effort-related drawback nevertheless draws a firm boundary against replacement:

"I ask the AI bot for clearer, gentler, more professional ways to express something, for example, professional ways to say that I do not like something. Now my emails look more professional, and it saves time. But relying on it most of the time also has drawbacks, especially in terms of cognitive effort, so it can assist but never replace humans." (P62)

Another participant associates teacher value with communicative domains that require embodied interaction and affective expression, while restricting AI use to functional rewriting:

"It helps me reformulate sentences and questions and write texts with fewer mistakes, and it can translate in a way that fits the field using technical terms. But it cannot replace a teacher when it comes to oral expression or expressing feelings, where the teacher's presence matters. For that reason, my own use would remain limited rather than total." (P05)

Across accounts, instructional value is linked less to "faster correction" than to how feedback is integrated into learning routines, how trust is calibrated for different tasks, and how human teaching is positioned as a safeguard for context-sensitive judgement and sustained academic work. Participants' rationales collectively situate the teacher as a calibration and alignment layer: teachers are associated with setting criteria for what counts as acceptable performance in a course, modelling how to evaluate feedback, and guiding learners to justify revisions rather than merely apply them. In the same accounts, AI is positioned as most useful when channelled toward low-stakes rehearsal and formative drafting (templates, reformulation options, practice prompts) and as least defensible when used as an unverified authority for sensitive, high-stakes, or context-dependent claims. This teacher-AI complementarity framing yields concrete classroom relevance that is already visible in the findings: learners describe stronger trust when verification and mediation practices are in place, and weaker trust when reliance becomes routinised or context drifts remain unchecked; in this sense, the "teacher role" is reported as a condition that shapes how AI feedback is taken up and whether it supports learning-to-learn practices.

5. Discussion

This Discussion interprets the findings as evidence that the value of AI-mediated feedback lies less in output generation itself than in the conditions under which learners inspect, verify, restrict, and integrate it into their work. Three claims organise the interpretation: first, motivation and autonomy are described as conditional on feedback-uptake routines rather than on access alone; second, trust calibration shapes how usefulness is judged across tasks; and third, teacher guidance remains central where academic stakes, contextual judgement, and meaning control are higher. On this basis, the Discussion advances an assessment-for-learning account in which educational value is associated with mediated uptake and selective trust rather than automated judgement.

5.1. TAM as a frame for condition-dependent adoption

TAM is useful here as a limited frame for understanding how participants describe ease and usefulness in relation to AI-mediated feedback. In the corpus, PEOU is commonly associated with rapid access, low-friction interaction, and support for routine drafting, quick checks, and practice exchanges. PU, by contrast, is described as more conditional: it aligns more strongly with low-stakes tasks and repeated iteration, and less strongly where tasks require sustained context-setting, meaning preservation, evidentiary caution, or careful stance and register control. Read in this way, TAM helps organise adoption-related judgements, but those judgements remain closely tied to task stakes and trust conditions rather than functioning as abstract attitudes.

5.2. Limits of TAM

TAM remains insufficient on its own for interpreting the present findings. Participants' accounts show that continued use does not imply unqualified acceptance, since uptake is often accompanied by verification, selective avoidance, manual editing, and explicit boundary-setting. The findings also indicate that access constraints and pragmatic-cultural fit shape what learners can do with AI in practice, not merely how positively they view it. For this reason, TAM is retained as an organising frame for adoption-related judgements, but interpretation of learning quality depends on additional concepts drawn from assessment-for-learning, especially feedback uptake, criteria alignment, and calibrated trust.

5.3. Learning as multidimensional performance refinement rather than a single metric

Participants' reports depict learning through a set of performance dimensions and task types rather than through a single indicator. The most salient targets cluster around accuracy (grammar and mechanics), fluency (naturalness and flow), controlled complexity (organisation and lexical choice), and genre/register appropriacy. These targets are articulated through writing and revision work, reformulation/translation tasks, and, for a subset, oral rehearsal and interaction practice. This task-based framing aligns with work describing chatbot use as associated with self-paced practice opportunities and reduced communication anxiety, while also relating to questions of how learning outcomes are defined and evaluated [3,1].

A tension is visible across accounts: convenience and speed are associated with motivation and persistence, yet fluency-oriented reformulation is also associated with concerns about meaning drift, stylistic uniformity, and reduced cognitive effort. This tension suggests that perceived benefit relates to domain and task type, and that fluent output is not uniformly treated as equivalent to improved communicative competence.

5.4. The learning journey as a feedback-use cycle

The findings support a process account in which AI feedback is treated as input to revision decisions rather than as a terminal judgement. Participants frequently describe a recurring feedback-use cycle: task specification → output inspection for fit (tone, coherence, lexical appropriacy) → diagnosis of weaknesses (genericness, meaning drift, factual uncertainty) → prompt revision and constraint tightening → adoption, manual editing, verification, or rejection. In a cross-sectional qualitative design, this cycle provides a defensible representation of learning as cumulative practice in self-regulation and feedback literacy.

Fig. 1 summarises this reported feedback-use cycle as a sequence of learner actions (specification, inspection, diagnosis, prompt refinement, and uptake decisions).

Note. Conceptual summary of the reported feedback-use cycle; not a frequency display or fixed sequence.

Fig. 1 clarifies the main claim of this section: the value of AI feedback lies in the cycle of inspection, diagnosis, revision, and uptake that follows generation, not in output delivery alone. It also sharpens the reading of dependency narratives. Statements about "addiction," reduced effort, or weakened searching habits are treated as perceived risks linked to routinised uptake and reduced verification rather than as direct evidence of altered attainment, aligning with related work [21,4,13].

5.5. For whom and under what conditions: learner profiles and moderation mechanisms

Conditional interpretation is supported through two analytic steps. First, the uptake profiles identified in the Findings (genre-polish/efficiency users, oral practice users, verification-oriented users, boundary-setters) can be read as strategy clusters associated with different configurations of benefit and risk. These are not treated as stable learner types; they index patterned ways of integrating feedback into learning practice.

Fig. 2 summarises the two moderation pathways described in participants' accounts and clarifies the article's conditional reading of usefulness. Access constraints and pragmatic-cultural fit do not simply qualify convenience; they shape the tasks for which AI feedback is treated as usable and trustworthy. Paywalls and unstable connectivity are associated with reduced iteration and a narrowed task range, while context-sensitive tasks are more often linked to verification or

avoidance. In TAM terms, PEOU may remain relatively high even where PU for extended academic tasks is lower [10,5,20].

Note. Conceptual summary of how access constraints and pragmatic-cultural fit shape task range, verification, and perceived usefulness.

Pragmatic-cultural fit operates as a trust moderator through task selection and verification. When tasks are culturally neutral, trust is more often described as sufficient for formative support. When tasks require contextual judgement (religious, historical, political topics; stance and register in locally sensitive discourse), verification and avoidance become more common. In these accounts, fluency is repeatedly distinguished from reliability and pragmatic appropriacy, which aligns with a multidimensional construct of learning.

5.6. Standards anchoring, teacher comparison, and safeguards for validity, fairness, and bias risk

The present evidence base is qualitative self-report and does not provide empirical validation of automated scoring, calibration stability, or bias measurement. Standards are therefore used as an interpretive and pedagogical anchor rather than as a measurement claim. A CEFR-informed lens can situate reported targets within widely recognised domains, grammatical accuracy, lexical range, coherence/cohesion, and sociolinguistic/pragmatic appropriacy, without assigning CEFR levels [2]. Within this frame, the findings align with a conditional interpretation: AI support is more consistently described as usable for sentence-level and drafting-adjacent work (accuracy and phrasing) and more contested where pragmatic appropriacy, contextual judgement, and meaning preservation are central.

Teacher comparison enters the corpus through complementarity claims, which align with research indicating that student and teacher acceptance of AI-assisted assessment can diverge and can be shaped by differing validity expectations [19]. The findings support a safeguards model aligned with assessment-for-learning practice:

1. Criterion anchoring: teacher-provided rubrics that separate accuracy, coherence, lexical appropriacy, and register, aligned with course outcomes.
2. Selective human checkpoints: teacher review of a subset of AI-assisted drafts framed as calibration and alignment rather than surveillance.

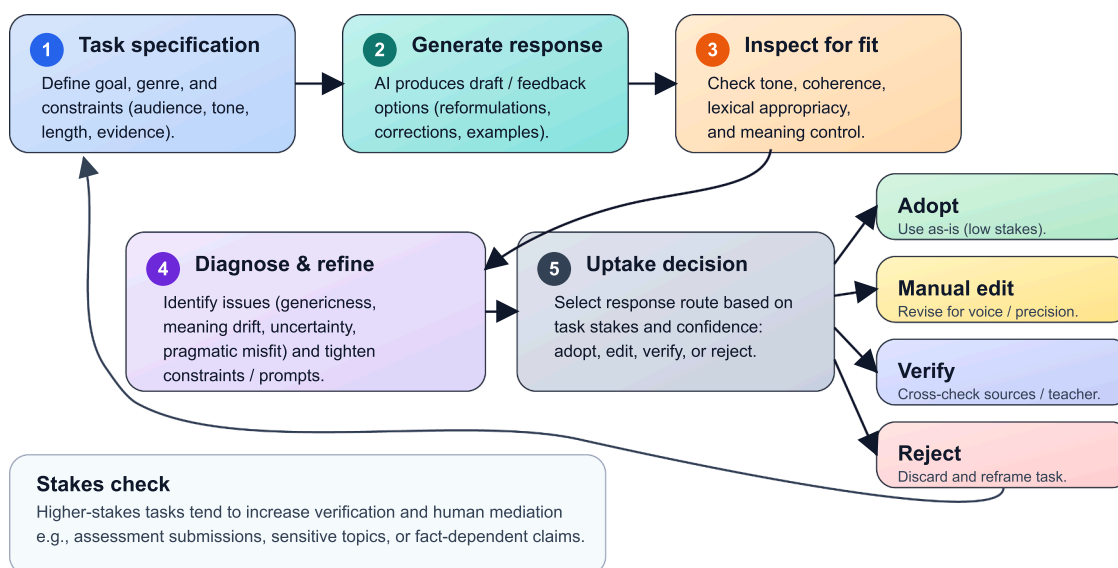


Fig. 1. Feedback-use cycle reported in participants' accounts of AI-mediated English learning.

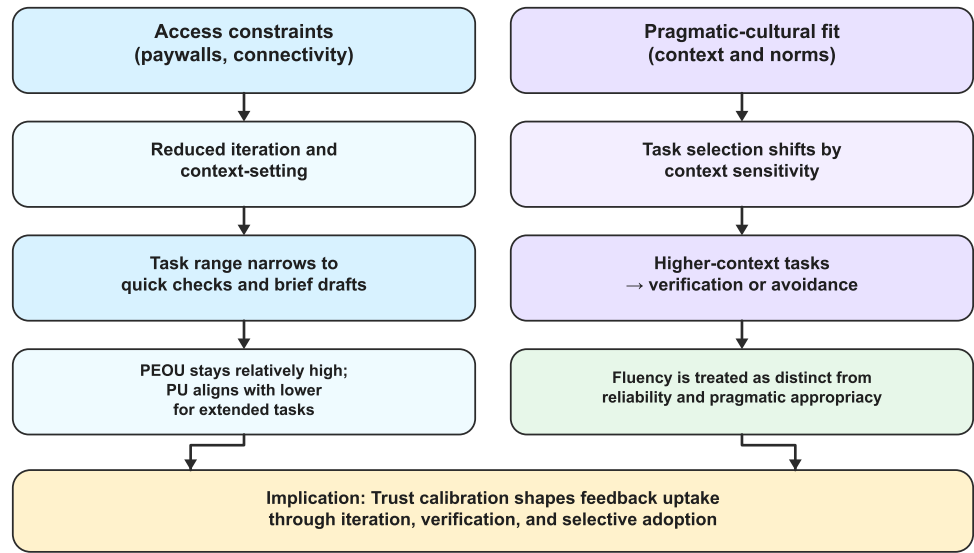


Fig. 2. Access constraints and pragmatic-cultural fit as moderators of usefulness and trust in AI-mediated feedback use.

- 3. Justification routines: learners provide brief rationales for adopted and rejected AI suggestions, including at least one rejected suggestion and the reason for rejection.
- 4. Verification expectations for consequential tasks: explicit rules for source checking and teacher/peer consultation when claims are factual, sensitive, or high-stakes.

These safeguards do not presume reliability of AI judgement; they specify conditions under which trustworthiness is more plausibly supported through mediation, verification, and criteria alignment.

5.7. Targeted pedagogical design: from general recommendations to implementable routines

Three implementable routines follow from the findings. First, a guided revision protocol can require the initial draft, prompt, AI output, revised draft, and a short audit of adopted and rejected changes, so evaluation attends to justification and revision quality rather than tool use. Second, error-analysis tasks can ask learners to identify meaning drift, factual uncertainty, or contextual misfit in an AI output and annotate it through rubric domains such as accuracy, coherence, register, and pragmatic appropriacy. Third, a culturally grounded prompt library can organise common Moroccan academic and professional

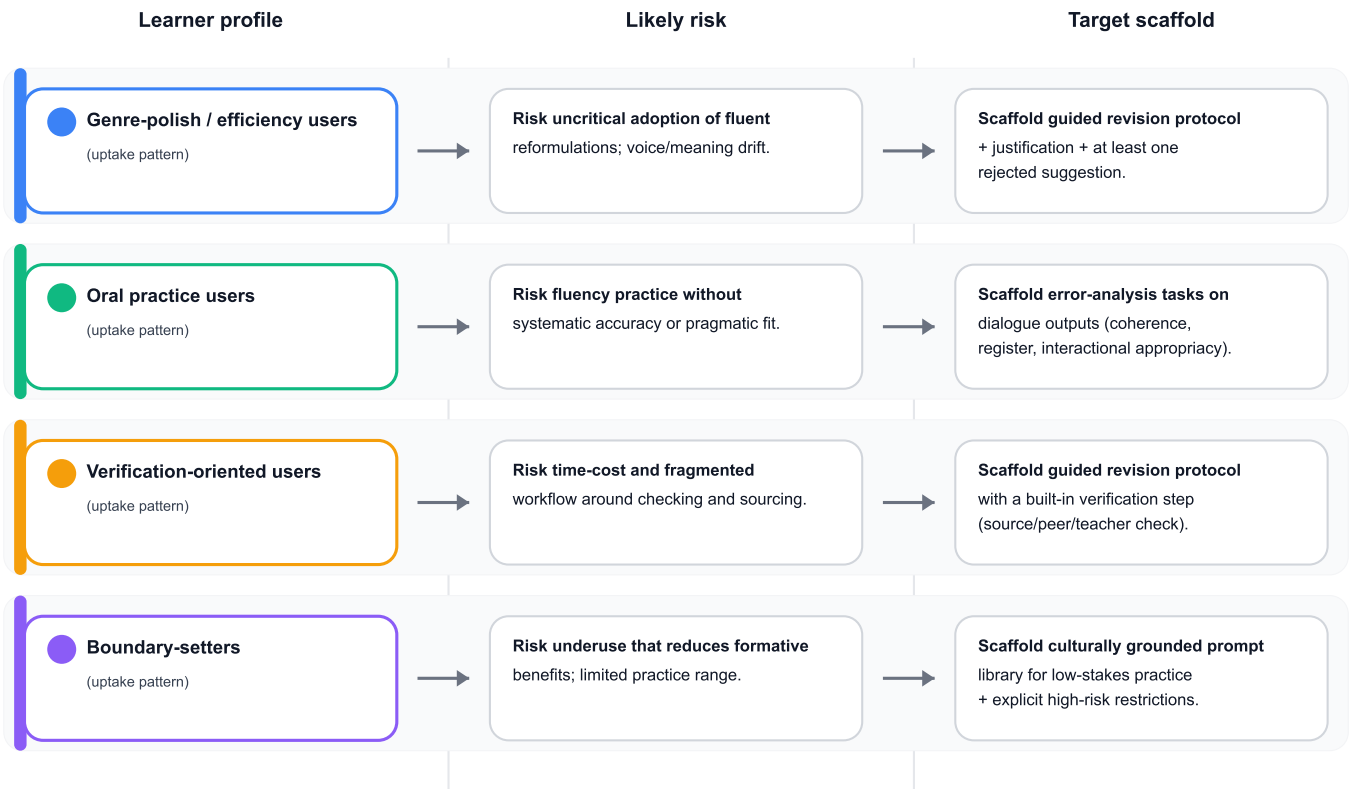


Fig. 3. Learner uptake profiles, likely risks, and targeted scaffolds for AI-mediated feedback use.

tasks, with verification tags for sensitive domains.

5.8. Linking learner profiles to scaffolds

Fig. 3 summarises the profile-to-scaffold logic by linking each uptake profile to a characteristic risk pattern and a targeted instructional response. The figure clarifies conditional interpretation by showing how the same AI tool use may relate to different learning consequences depending on uptake routines and task stakes. In this framing, scaffolds support feedback literacy and criterion alignment rather than serving as general advice to use AI cautiously.

Note. Conceptual mapping of uptake profiles, associated risks, and possible scaffolds; profiles are analytic rather than mutually exclusive.

The four uptake profiles suggest differentiated scaffolding. For genre-polish and efficiency-oriented users, the likely risk relates to uncritical adoption of fluent reformulations; the guided revision protocol conditions uptake on justification and documented rejection decisions. For oral practice users, the likely risk relates to extended fluency practice with limited attention to accuracy or pragmatic fit; error-analysis tasks applied to dialogue outputs can redirect attention to coherence, register, and interactional appropriacy. For verification-oriented users, the likely risk relates to time-cost and workflow fragmentation; embedding a structured verification step within the guided revision protocol preserves evaluative stance while maintaining feasibility. For boundary-setters, the likely risk relates to underuse that restricts formative benefit; the culturally grounded prompt library supports low-stakes, context-aware practice while retaining explicit restrictions for high-risk domains.

5.9. Limitations and transferability

Several limitations condition interpretation. The dataset derives from a single Moroccan university cohort and from a convenience sample of students who had already used AI tools for English learning; the findings therefore reflect self-selected participant accounts rather than a broad institutional profile. The design is cross-sectional and based on self-reported open-ended survey responses, so the analysis concerns reported practices, perceived risks, and trust conditions rather than demonstrated learning gains or longitudinal change. No demographic or proficiency variables were collected, which restricts subgroup interpretation, and multilingual responses were translated into English for analysis, a procedure that preserves propositional meaning but may relate to minor shifts in tone. Transferability is therefore analytic rather than statistical and is strongest for comparable higher-education EFL contexts shaped by resource constraints, academic writing demands, and context-sensitive use of AI-mediated feedback.

5.10. Boundary conditions and future work aligned with assessment expectations

The evidence supports claims about learners' perceptions, uptake practices, and trust conditions; it does not support claims about measured learning gains or declines. Dependency statements are therefore treated as perceived risks associated with self-regulation rather than as objective evidence of altered learning quality. Future work aligned with assessment expectations could combine learner accounts with (a) longitudinal writing samples, (b) teacher ratings using rubric-guided criteria that reference CEFR domains, (c) comparison of AI- and teacher-generated feedback on the same texts to examine convergence and divergence, and (d) a focused audit of culturally sensitive prompts to examine reliability and fairness under local conditions. Such designs would permit stronger evidence about growth over time, domain-specific performance trajectories, and the extent to which feedback supports learning-to-learn practices in situated classrooms.

6. Conclusion

In this cohort, Moroccan graduate students describe AI-mediated feedback as educationally useful not in general or unconditional terms, but under specific uptake and trust conditions. Reported usefulness aligns most clearly with task-situated language work such as drafting, revision, reformulation, and rehearsal, while motivation and autonomy are described as stronger when learners inspect outputs, refine prompts, verify claims, and restrict use where stakes are higher. Interpreted through TAM, PEOU remains high for routine tasks, whereas PU is more clearly condition-dependent and shaped by access, contextual fit, and confidence in the output. Teachers therefore remain central not as mere providers of correction, but as calibration and alignment agents whose role is most salient where judgement, criteria, and context-sensitive meaning matter most.

For practice and policy, the findings suggest that AI is most defensibly integrated as an assessment-for-learning scaffold with explicit safeguards. Three routines follow directly from the uptake patterns described: (1) a guided revision protocol requiring the initial draft, prompt, AI output, revised draft, and a brief justification of adopted and rejected changes; (2) structured error-analysis activities in which learners identify meaning drift, factual uncertainty, or pragmatic misfit and annotate these using rubric domains (accuracy, coherence, lexical appropriacy, register); and (3) a locally grounded prompt library for common Moroccan academic and professional tasks, with "verification required" tags for sensitive topics and consequential claims. At classroom level, the implication is that feedback value relates to mediation and evaluation (verification checkpoints, teacher/peer review of a subset of AI-assisted drafts, and criteria alignment via course rubrics that can be read in relation to CEFR-associated domains) rather than to unqualified adoption of fluent text. At institutional level, access conditions remain central: stable connectivity, transparent guidance for free-tier constraints, and teacher development oriented to evaluation and contextual judgement are likely to support more equitable uptake.

These conclusions are bounded by the study design. Evidence is cross-sectional and self-reported, and multilingual responses were translated into English for reporting; therefore, claims concern described practices, perceived risks, and trust conditions rather than measured gains, calibration stability, or bias estimation in automated scoring. Future work aligned with assessment expectations could combine learner accounts with longitudinal writing samples, teacher ratings using rubric-guided criteria, and direct comparison of AI- and teacher-generated feedback on the same texts, alongside a focused audit of context-sensitive prompts to examine reliability and fairness under local conditions.

Funding

This research received no external funding.

CRediT authorship contribution statement

Nouh Alaoui Mhamdi: Writing – original draft, Data curation, Conceptualization.

Declaration of competing interest

There are no financial or personal relationships that could inappropriately influence or bias the content of the paper.

Acknowledgements

The author wishes to thank colleagues at the Faculty of Letters and Human Sciences, Dhar Mahraz, for their feedback during the design of this study, and the participating graduate students for their time and insights.

Appendix A. Verbatim open-ended survey prompts (French)

Instrument. The following seven prompts elicited the qualitative data analysed in the study; they are reproduced verbatim in French. Responses were permitted in French, Arabic, or English; non-English quotations in the Findings are translated.

1. Comment utilisez-vous des outils d'intelligence artificielle, comme ChatGPT, pour apprendre ou améliorer votre anglais ? Pouvez-vous donner des exemples concrets ?
2. En quoi l'utilisation d'outils comme ChatGPT a-t-elle changé votre manière de rédiger des textes ou de communiquer en anglais dans un cadre académique ou professionnel ?
3. Quels sont, selon vous, les avantages de ChatGPT ou d'autres outils d'intelligence artificielle dans l'apprentissage de l'anglais ? Et quels sont leurs inconvénients ?
4. Avez-vous rencontré des défis ou des obstacles en utilisant ChatGPT pour apprendre l'anglais ? Si oui, lesquels, et comment les avez-vous surmontés ?
5. Pensez-vous que ChatGPT peut remplacer un enseignant d'anglais ? Pourquoi ou pourquoi pas ?
6. Comment voyez-vous l'utilisation d'outils comme ChatGPT dans votre apprentissage futur de l'anglais ou dans votre carrière professionnelle ?
7. Y a-t-il autre chose que vous aimeriez ajouter concernant votre expérience avec ChatGPT ou d'autres outils d'intelligence artificielle dans l'apprentissage ou l'utilisation de l'anglais ?

References

- [1] Jeon J. Exploring AI chatbot affordances in the EFL classroom: young learners' experiences and perspectives. *Comput Assist Lang Learn* 2024;37(1–2):1–26. <https://doi.org/10.1080/09588221.2021.2021241>.
- [2] Law L. Application of generative artificial intelligence (GenAI) in language teaching and learning: a scoping literature review. *Comput Educ Open* 2024;6: 100174. <https://doi.org/10.1016/j.caeo.2024.100174>.
- [3] Hong JC, Lin CH, Juh CC. Using a chatbot to learn English via charades: the correlates between social presence, hedonic value, perceived value, and learning outcome. *Interact Learn Environ* 2024;32(10):6590–606. <https://doi.org/10.1080/10494820.2023.2273485>.
- [4] Salam U. The integration of ChatGPT in English for foreign language course: elevating AI writing assistant acceptance. *Comput Sch* 2024;1–21. <https://doi.org/10.1080/07380569.2024.2446239>.
- [5] Kshetri N. Linguistic challenges in generative artificial intelligence: implications for low-resource languages in the developing world. *J Glob Inf Technol Manag* 2024;27(2):95–9. <https://doi.org/10.1080/1097198X.2024.2341496>.
- [6] Liu G, Ma C. Measuring EFL learners' use of ChatGPT in informal digital learning of English based on the technology acceptance model. *Innov Lang Learn Teach* 2024; 18(2):125–38. <https://doi.org/10.1080/17501229.2023.2240316>.
- [7] Hadinejad N, Sperling K, McGrath C. Generative AI chatbots in higher education: student experiences and perceived ethical challenges. *Comput Educ Open* 2025;9: 100311. <https://doi.org/10.1016/j.caeo.2025.100311>.
- [8] Zou S, Guo K, Wang J, Liu Y. Investigating students' Uptake of teacher- and ChatGPT-generated feedback in efl writing: a comparison study. *Comput Assist Lang Learn* 2025;1–30. <https://doi.org/10.1080/09588221.2024.2447279>.
- [9] Alaoui Mhamdi N. Questioning the narrative of generative AI and the metaverse in education: a mixed-methods study of Moroccan teachers' Perceptions and behavioural intentions (Eds.). In: Vallejo RGá, Moukhliis G, Schaeffer E, Paliktzoglou I, editors. *Proceedings of the second international symposium on generative AI and education (ISGAIE'2025)* 262. Springer; 2025. p. 471–86. https://doi.org/10.1007/978-3-031-98476-1_37.
- [10] Bouaziz K, Eddouada S. AI-assisted academic writing and publishing in Morocco: doctoral students' Ethical decision-making accounts. *Arab World Engl J* 2026;(3): 268–86. <https://doi.org/10.24093/awej/A3.17>.
- [11] Davis FD. Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Q* 1989;13(3):319. <https://doi.org/10.2307/249008>.
- [12] Alfaleh M, Albais HH, Mohamed AM. EFL learners' Motivation and engagement with oral reading activities and AI tools: exploring the effect of gender, study level and English language proficiency. *Cogent Educ* 2025;12(1):2542395. <https://doi.org/10.1080/2331186X.2025.2542395>.
- [13] Shi L, Muhammad Umer A, Shi Y. Utilizing AI models to optimize blended teaching effectiveness in college-level English education. *Cogent Educ* 2023;10(2):2282804. <https://doi.org/10.1080/2331186X.2023.2282804>.
- [14] Zhang Z, Aubrey S, Huang X, Chiu TKF. The role of generative AI and hybrid feedback in improving L2 writing skills: a comparative study. *Innov Lang Learn Teach* 2025;1–19. <https://doi.org/10.1080/17501229.2025.2503890>.
- [15] Mjahad R, Boukranaa A, ElKarfa A, Sandy K. The integration of flipped learning in Moroccan higher education: benefits and challenges. *Cogent Educ* 2025;12(1): 2546573. <https://doi.org/10.1080/2331186X.2025.2546573>.
- [16] Belda-Medina J, Kokošková V. ChatGPT for language learning: assessing teacher candidates' skills and perceptions using the technology acceptance model (TAM). *Innov Lang Learn Teach* 2024;1–16. <https://doi.org/10.1080/17501229.2024.2435900>.
- [17] Moussa A, Belhiah H. Beyond syntax: exploring Moroccan undergraduate with AI-assisted writing. *Arab World Engl J* 2024;1(1):138–55. <https://doi.org/10.24093/awej/ChatGPT.9>.
- [18] Amaning N. Qualitative analysis on technology acceptance model of accounting faculty perceptions of virtual accounting coursework. *Cogent Educ* 2024;11(1): 2331345. <https://doi.org/10.1080/2331186X.2024.2331345>.
- [19] Van Den Berg S, Papadopoulos PM. Summative assessment with artificial intelligence: qualitative analysis and comparison of technology acceptance in student and teacher populations. *Innov Educ Teach Int* 2024;1–16. <https://doi.org/10.1080/14703297.2024.2436613>.
- [20] O'Dea X. Generative AI: is it a paradigm shift for higher education? *Stud High Educ* 2024;49(5):811–6. <https://doi.org/10.1080/03075079.2024.2332944>.
- [21] Pallant JL, Blijlevens J, Campbell A, Jopp R. Mastering knowledge: the impact of generative AI on student learning outcomes. *Stud High Educ* 2025;1–22. <https://doi.org/10.1080/03075079.2025.2487570>.