



From AI Dependence to AI Literacy: Addressing Neural Machine Translation Biases in English–Arabic Digital Media Translation

Brahim Machaal
Ibn Zohr University, Morocco

Abstract

This empirical study examines how 45 Moroccan undergraduate students in Applied Foreign Languages translated English digital media content into Modern Standard Arabic (MSA). Analysing 360 translations across eight genres; including social media posts, streaming content descriptions, and mobile application interfaces, two independent raters applied a seven-category taxonomy adapted from Multidimensional Quality Metrics (MQM) frameworks, with strong interrater agreement (Cohen's $\kappa = 0.84$). Findings reveal systematic error patterns consistent with documented neural machine translation (NMT) biases: over-formalisation (73% of translations), inappropriate literal transfer of cultural references (68%), and platform constraint violations (62%). Drawing on an extended PACTE translation competence model, the study argues that students lack the critical AI literacy required to evaluate machine translation output; particularly regarding register, cultural adaptation, and platform-specific conventions, despite ready access to tools such as Google Translate and ChatGPT. Follow-up interviews indicate that students evaluated AI output primarily on surface grammatical correctness, demonstrating insufficient strategic competence to identify functionally inappropriate output. A pedagogical framework is proposed that embeds NMT bias awareness, post-editing skills, and AI tool evaluation as assessable competency milestones. The results call for translation programmes to reposition critical AI literacy from a peripheral concern to a core curricular component structured around human–AI complementarity.

Keywords: neural machine translation, AI translation tools, translation pedagogy, translation error analysis, English–Arabic translation, digital media localisation

1. Introduction

Neural machine translation has fundamentally transformed professional translation practice, yet translation pedagogy has been slow to address how students actually use these technologies and the systematic errors that result from uncritical AI adoption. Digital media

translation—encompassing social media posts, streaming service content, mobile application interfaces, and influencer marketing copy—presents particularly acute challenges because AI systems trained predominantly on formal corpora systematically fail to capture the informal registers, cultural nuance, and platform-specific conventions essential to this domain (Castilho et al., 2022; Läubli et al., 2020).

For Moroccan translation students, this technological reality creates a pedagogical paradox. AI translation tools are ubiquitous and produce superficially fluent Arabic output, yet those same systems exhibit systematic biases that make their output functionally inappropriate for digital media contexts despite grammatical correctness (Sajjad et al., 2020; Nagoudi et al., 2022). It has been observed that Google Translate and ChatGPT consistently over-formalise casual English content when translating to Arabic, literally transfer cultural references that require creative adaptation, and disregard platform constraints such as character limits or interface conventions—which are precisely the competencies that digital media translators must master.

Current translation pedagogy inadequately addresses this reality. In fact, most programmes treat AI tools peripherally, focusing on traditional translation theory while students routinely use NMT systems without critical evaluation frameworks (Mellinger & Hanson, 2020; Kenny, 2022). This gap is particularly problematic given that AI-assisted translation is now standard professional practice, requiring not avoidance but strategic, critical engagement.

This study documents how Moroccan translation students use AI translation tools, the systematic errors arising from uncritical adoption, and the pedagogical interventions necessary to develop critical AI literacy alongside traditional translation competencies. Three research questions guide the investigation:

- What patterns of AI tool usage and AI-mediated errors appear in student translations of English digital media content into Arabic?
- What systematic limitations of NMT systems for English–Arabic digital media translation do student error patterns reveal?
- How must translation pedagogy evolve to develop critical AI literacy and strategic technology engagement as core competencies?

2. Literature Review

2.1 Translation Competence in the Digital Era

Translation competence research provides the foundational frameworks for understanding translator development. The PACTE group (2003, 2005) identifies six interrelated sub-competencies: bilingual, extralinguistic, knowledge about translation, instrumental, psycho-physiological, and strategic, with strategic competence coordinating the others. Crucially, this model underscores that translation competence exceeds bilingual ability, requiring specialised knowledge and metacognitive processes. Subsequent empirical validation of the PACTE model across multiple learner populations has confirmed its applicability to diverse translation contexts (Hurtado Albir, 2017).

Göpferich (2009) extended the PACTE framework and theorised competence acquisition stages, proposing that novice translators initially rely heavily on linguistic competencies before developing strategic competencies through sustained practice. Her longitudinal studies of translation students revealed predictable developmental trajectories, with error patterns serving as diagnostic indicators of competence gaps. This insight suggests that systematic

error analysis can yield targeted information for pedagogical interventions addressing specific developmental needs.

Recent revisions of competence models foreground technological mediation as central to contemporary practice (Massey & Ehrensberger-Dow, 2017; Kiraly, 2015). Mellinger and Hanson (2020) argue that translation competence must now encompass the critical evaluation of machine translation output, an understanding of NMT limitations, and the ability to leverage AI tools strategically. For digital media translation specifically, Jiménez-Crespo (2013) argues that traditional models require expansion to include technological literacy, multimodal awareness, and knowledge of user experience principles. Empirical studies of professional digital media translators confirm that successful practitioners integrate technical, linguistic, and user-centred design competencies in ways not captured by earlier models (Jiménez-Crespo, 2013; Kenny, 2022).

2.2 Error Analysis in Translation Pedagogy

Error analysis serves both as an assessment methodology and a pedagogical tool (James, 2013). Hansen's (2010) classification system distinguishes competence errors (knowledge gaps), performance errors (lapses despite adequate knowledge), and pragmatic errors (failures to achieve communicative purpose). This tripartite framework enables educators to distinguish between errors requiring new instruction and those addressable through practice and metacognitive training. Kenny (2022) extends these categories to machine translation contexts, proposing additional categories for MT-induced errors (problems originating in automated output), post-editing under-correction (failure to fix MT errors), and post-editing over-correction (unnecessary modifications).

For Arabic translation specifically, Husni and Newman (2015) document persistent challenges students face; namely, managing cultural adaptation decisions across diverse Arab audiences, handling terminology gaps where Arabic equivalents are unstable or contested, and negotiating between source culture preservation and target culture adaptation. Their analysis reveals that Arabic learners particularly struggle with register appropriateness, frequently producing overly formal translations regardless of source text register. These patterns confirmed claims by other studies of Arabic translation pedagogy (Farghal & Shunnaq, 1999).

2.3 Neural Machine Translation: Systematic Limitations for Digital Media

Understanding NMT limitations is essential for developing critical AI literacy in translation students. Koehn and Knowles (2017) identified six fundamental NMT challenges directly relevant to digital media translation: domain mismatch, low-resource language pair difficulties, rare word handling (neologisms and internet slang), long sentence processing, inappropriate lexical choices, and beam search errors that select fluent but inaccurate translations.

For English–Arabic translation, Bouamor et al. (2018) demonstrated that NMT systems perform optimally on formal news text, with quality degrading significantly for informal social media content, which is precisely the register dominating contemporary digital communication. This performance gap stems largely from training data composition. Indeed, Arabic NMT corpora derive predominantly from news agencies, government documents, and formal literary texts, with informal digital Arabic considerably underrepresented in the pre-transformer era (Saadane & Habash, 2015; Darwish et al., 2021).

Research on Arabic NMT has documented the critical systematic tendency that Arabic NMT models trained on formal corpora over-formalise output regardless of source text register.

Sajjad et al. (2020), whose AraBench benchmark evaluation highlights performance degradation on dialectal and informal Arabic, illustrates the broader training data mismatch: systems optimised on formal written Arabic produce Modern Standard Arabic with formal vocabulary and grammatical constructions inappropriate for casual communicative contexts. This is not an occasional error but a systematic pattern rooted in training data composition, confirmed across multiple evaluation studies (Bouamor et al., 2018; Nagoudi et al., 2022).

Equally, evaluations of large language models (LLMs) for Arabic translation reveal a more nuanced picture. Hendy et al. (2023) found that GPT-4 outperforms Google Translate on Arabic across multiple quality metrics, and Jiao et al. (2023) confirmed competitive multilingual performance of ChatGPT. Nevertheless, both studies note that even LLMs exhibit register inconsistencies and cultural reference literalism in informal text domains, suggesting the limitation is not absent in LLMs but reduced compared to earlier NMT systems. For translation pedagogy, the key insight that remains consistent across system generations is that grammatically flawless output can be functionally inappropriate when register, cultural adaptation, or platform conventions are not addressed.

2.4 Cultural Adaptation in Arabic Translation

Cultural adaptation presents distinct challenges in Arabic translation. Dickins et al. (2017) propose strategies ranging from exoticism (source culture retention with explanatory additions) through calque and cultural transplantation (target culture substitution) to communicative translation (functional equivalence prioritising target audience comprehension). These frameworks prove essential for digital media, where audiences expect both international brand consistency and local cultural relevance. Empirical studies of professional Arabic localisers reveal that successful practitioners develop sophisticated decision-making heuristics balancing brand identity preservation with cultural appropriateness (Faiq, 2004).

Research on Arabic digital media localisation emphasises the importance of understanding pan-Arab versus regional cultural references, given the linguistic and cultural diversity across Arabic-speaking contexts. Studies in Arabic sociolinguistics and translation practice confirm that what resonates culturally in one regional Arabic variety may not transfer effectively to audiences in other Arabic-speaking countries, requiring translators to make deliberate decisions about target audience scope and cultural specificity (Dickins et al., 2017; Husni & Newman, 2015).

3. Theoretical Framework

This study draws on three interconnected frameworks adapted for AI-mediated digital translation contexts.

3.1 PACTE Competence Model

The PACTE (2003, 2005), as extended by Massey and Ehrensberger-Dow (2017), provides the primary framework for error analysis through sub-competencies. The extended version foregrounds instrumental competence (AI tools, post-editing platforms, content management systems) and strategic competence (coordinating human judgement with technological capabilities), enabling systematic analysis of whether errors stem from linguistic knowledge gaps, technological limitations, or strategic decision-making failures.

3.2 Hansen's Error Typology

Hansen's (2010) typology guides error classification, distinguishing between competence, performance, and pragmatic errors. Kenny's (2022) extension adds MT-induced errors, post-editing under-correction (failure to identify and correct MT errors), and post-editing over-correction (unnecessary or harmful modifications to adequate MT output) as categories directly relevant to AI-mediated translation. This combined framework allows identification not only of what errors occur but of their underlying cognitive and technological sources, including cases where students modify machine output that was contextually appropriate and in doing so introduce new errors.

3.3 Nord's Functionalist Approach

Nord (2018) Skopos theory provides the framework for evaluating cultural adaptation and platform appropriateness, positing that translation adequacy should be evaluated against communicative purpose rather than source text fidelity alone. This theoretical lens is essential for digital media translation, where communicative effectiveness consistently takes precedence over literal correspondence.

4. Methodology

4.1 Participants and Context

All 45 Moroccan undergraduate students in the final semester (Semester 6) of a Bachelor's degree in Applied Foreign Language Studies at Ibn Zohr University, Morocco participated in this study, conducted between October 2024 and March 2025. All participants were native Arabic speakers with extensive formal education in Modern Standard Arabic. Students held intermediate to upper-intermediate English proficiency (CEFR B1–B2). Although participants had completed an introductory translation course in Semester 4 covering translation tools and methods, they were concurrently enrolled in their first systematic advanced translation course offering structured instruction in translation theory, quality assessment, and AI-assisted translation. Participants therefore had prior informal experience with NMT tools but without structured training in critical evaluation or post-editing. Informed consent was obtained from all participants, who were free to withdraw at any time. All students completed and submitted their translations. Participants are identified using generic labels (Student A, Student B, etc.) to protect privacy.

4.2 Data Collection

Data derived from eight authentic English digital media texts produced between 2022–2024, representing diverse contemporary genres: a product announcement tweet, an Instagram fashion influencer caption, a YouTube cooking channel description, a Netflix teen drama synopsis, a mobile fitness application interface, a Facebook environmental NGO post, a TikTok spoken transcript, and a remote work blog excerpt (see Appendix A for full source texts). Selection criteria included: (1) authenticity (real content from actual platforms); (2) contemporary relevance; (3) linguistic diversity, spanning varied formality levels and internet-specific language; (4) genre representativeness across major digital media categories; and (5) graduated difficulty.

Students received instructions for a simulated professional assignment specifying: target audience (general Arabic-speaking users across Arab countries); purpose (enabling Arabic audiences to access and engage with content as source audiences do); quality expectations (professional publication quality appropriate to each platform); and a two-week completion

timeline. Students were permitted to use any resources typically available to professional digital content translators, including dictionaries, terminology databases, search engines, and AI tools. This approach mirrors realistic professional conditions, where AI tools are standard and restrictions on their use are neither enforceable nor professionally relevant. Submissions were evaluated on final quality and platform appropriateness, consistent with professional evaluation criteria that prioritise deliverable quality over production method.

4.3 Process Limitations

A methodological constraint requires explicit acknowledgement. Without keystroke logging, screen recording, or think-aloud protocols (cf. Englund Dimitrova, 2005; Göpferich et al., 2008), it is not possible to determine definitively which errors originated in human translation decisions versus AI-generated output accepted without modification. Think-aloud protocols were infeasible because translation tasks occurred over two weeks in students' independent settings, preventing real-time researcher monitoring and recording. While follow-up interviews with a representative sample (Section 4.5) provide evidence of AI tool usage patterns, individual error attribution remains probabilistic rather than certain. AI influence is inferred through pattern matching: when student errors systematically mirror documented NMT biases—over-formalisation, literal cultural reference translation, platform constraint violations—AI mediation is a plausible explanation. Alternative explanations, including shared pedagogical influences or common knowledge gaps, cannot be entirely excluded. All causal claims are accordingly qualified throughout.

4.4 Analysis Procedure

Two experienced evaluators—one specialising in translation pedagogy (10 years' experience) and one in professional Arabic digital media translation (6 years' experience)—independently assessed all 360 translations (45 students $\times$ 8 texts). Prior to rating, the evaluators completed a calibration session in which they jointly rated 20 translations not included in the main dataset, discussed discrepancies, and reached consensus on borderline cases; this process ensured shared interpretive criteria before independent assessment commenced. Where a single translation contained multiple instances of the same error category, each instance was counted separately to capture the full distribution of error types; for cross-category overlaps, raters were instructed to assign each error to its primary functional category. Errors were classified using a taxonomy adapted from Multidimensional Quality Metrics (MQM) frameworks (Lommel et al., 2014) for digital media translation contexts. The taxonomy comprised seven categories:

- (1) cultural adaptation errors;
- (2) lexical errors (terminology inconsistency, neologism mishandling, false friends, register-inappropriate word choice);
- (3) technical and platform-specific errors (character limit violations, hashtag mishandling, interface terminology);
- (4) syntactic errors;
- (5) pragmatic errors (tone mismatch, register failure, communicative purpose failure);
- (6) omission and addition errors; and
- (7) orthographic and typographic errors. Inter-rater reliability, calculated using Cohen's kappa, was strong ($\kappa = 0.84$). Discrepancies were resolved through structured discussion.

Error frequencies were calculated across all translations to identify prevalent patterns.

4.5 Follow-up Interviews

Semi-structured interviews were conducted with 10 purposefully sampled students selected to represent the full performance range across the translation tasks: three high performers (top quartile), four mid-range performers, and three low performers (bottom quartile). Interviews of 45–60 minutes explored translation processes, resource and AI tool usage, awareness of digital Arabic conventions, and attitudes toward AI assistance. Data were analysed thematically following Braun and Clarke (2006). Eight of the ten interviewed students reported using AI tools, primarily Google Translate and ChatGPT, with varying degrees of critical engagement. None of the interviewed students reported using DeepL, even though it was freely available via its website.

5. Findings

Error frequency analysis revealed systematic patterns across the 360 translations. Cultural adaptation errors appeared in 73% of translations, followed by lexical errors (68%), technical and platform-specific errors (62%), syntactic errors (51%), pragmatic errors (44%), omission and addition errors (33%), and orthographic and typographic errors (28%). Error rates varied by genre, with Instagram captions showing the highest rate of cultural adaptation errors (89%) and fitness application interfaces showing the highest rate of technical errors (82%). Table 1 presents the full error distribution. Table 2 (Appendix B) provides the complete genre-specific breakdown.

Table 1. Error Distribution Across All Translations (N = 360)

Error Category	Frequency (%)	No. of Translations	Mean Errors per Translation
Cultural Adaptation	73	263	2.1
Lexical	68	245	1.9
Technical / Platform-Specific	62	223	1.7
Syntactic	51	184	1.4
Pragmatic	44	158	1.2
Omission / Addition	33	119	0.8
Orthographic / Typographic	28	101	0.6

5.1 Cultural Adaptation Errors (73%)

Cultural adaptation errors manifested in three sub-patterns within the 73% total: over-foreignization (accounting for 51% of cultural adaptation errors), pop culture reference mishandling (34%), and inappropriate domestication (28%). Note that these sub-percentages reflect the proportion of cultural adaptation errors falling into each pattern, not the proportion of all 360 translations; individual translations sometimes showed more than one sub-pattern.

Over-Foreignization Example 1 (Instagram):

“Giving major cottage core vibes with this thrifted find! Sustainable fashion is the new black, and I’m here for it 🍷”

Student translation:

تعطي أجواء كوتاج كور كبيرة مع هذا الاكتشاف المستعمل! الموضة المستدامة هي الأسود الجديد، وأنا هنا من أجلها.

Problems: The transliterated “كوتاج كور” (cottage core) remains incomprehensible to Arabic readers; the literally translated “الأسود الجديد” (the new black) loses the fashion idiom’s meaning; the word-for-word “أنا هنا من أجلها” produces awkward Arabic failing to convey enthusiastic endorsement.

Suggested correction:

أتبنى أسلوب الحياة الريفية البسيطة مع هذا الكنز المستعمل! الموضة المستدامة أصبحت الاتجاه الأكثر رواجاً، وأنا أؤيدها بشدة. 

Over-Foreignization Example 2 (TikTok):

Source: “POV: You’re a millennial trying to explain to Gen Z why we still use the laughing-crying emoji 😂 Not me having a whole existential crisis over this!”

Student translation:

وجهة نظر: أنت من جيل الألفية تحاول أن تشرح للجيل زد لماذا ما زلنا نستخدم رمز الضحك والبكاء. ليس أنا أعاني من أزمة وجودية كاملة حول هذا!

Problems: The literal “وجهة نظر” (point of view) for “POV” misses that on TikTok “POV” functions as a specific format convention; “ليس أنا أعاني” produces grammatically awkward Arabic losing the ironic tone.

Suggested correction:

تخيل: شخص من جيل الألفية - مواليد الثمانينات والتسعينات - يحاول شرح لجيل زد - مواليد الألفين - لماذا ما زلنا نستخدم إيموجي الضحك مع الدموع 😂 وأنا طبعاً لست في أزمة وجودية بسبب هذا الموضوع. بل أنا كذلك فعلاً!

Inappropriate Domestication Example (YouTube):

Source: “Join me every Tuesday for Taco Tuesday! We’ll explore authentic Mexican street food recipes that’ll transport you straight to Guadalajara.”

Student translation:

انضموا إلي كل ثلاثاء لثلاثاء الطاجين! سنستكشف وصفات الطعام المغربي الأصيلة التي ستنتقلكم مباشرة إلى مراكش.

Problems: Replacing “Taco Tuesday” with “ثلاثاء الطاجين” (Tajine Tuesday) and “Guadalajara” with “مراكش” (Marrakech) fundamentally misrepresents the channel content, misleading viewers.

Suggested correction:

انضموا إلي كل ثلاثاء في يوم التاكو المكسيكي الخاص! سنستكشف وصفات أصيلة من طعام الشارع المكسيكي ستجعلكم تشعررون وكأنكم في مدينة غوادالاخارا المكسيكية.

5.2 Lexical Errors (68%)

Lexical errors clustered in four areas: technology terminology inconsistency (42%), neologism handling difficulties (38%), false friends and polysemy (31%), and inappropriate formality in word choice (27%).

Technology Terminology Inconsistency (42%):

For “Push Notifications” in mobile app settings, student translations varied dramatically: التنبيهات (9 students), الإشعارات المنبقة (10 students), الإشعارات الفورية (11 students), إشعارات الدفع (3 students), الإشعارات التلقائية (4 students), إشعارات بوش (8 students), الإشعارات الفورية (3 students).

Problem: Extreme variability, sometimes even within single students’ work across texts, suggests limited exposure to standardisation efforts, inconsistent terminology resource consultation, and lack of parallel text examination.

Suggested standardisation: الإشعارات الفورية.

Rationale: Major technology companies including Apple (Arabic Human Interface Guidelines), Google (Android Arabic Language Resources), and Samsung (Galaxy UI Arabic Style Guide) have standardised on “الإشعارات الفورية” (instant notifications) for their Arabic interfaces; for example, Apple’s iOS Arabic localisation consistently uses this string in Settings > Notifications. Consistency in technology terminology is crucial for user experience—users navigating app settings need consistent terminology throughout to avoid confusion.

Neologism Handling – “Zoom Fatigue” (38%):

Source (blog post): “Remote work has blurred the boundaries between work and life, leading to ‘Zoom fatigue’ and difficulty in ‘switching off’ after work hours.”

Student translations for “Zoom fatigue”: الإرهاق من زوم (11 students), تعب زوم (13 students), إرهاق الاجتماعات عبر الإنترنت (8 students), التعب بسبب الاجتماعات الافتراضية (9 students), ظاهرة تعب زوم (4 students).

Problems: Most translated “fatigue” as simple تعب (tiredness) without capturing the specific psychological/medical connotation; no students demonstrated awareness that “Zoom fatigue” has become a semi-technical term; none added explanatory context; none consulted contemporary Arabic usage.

Suggested correction (Option 1):

إرهاق زوم – الإرهاق الناتج عن كثرة الاجتماعات المرئية.

Rationale: Arabic media increasingly use “إرهاق زوم” as semi-standardised equivalent; parenthetical explanation ensures comprehension for readers unfamiliar with the term while introducing professional vocabulary.

Suggested correction (Option 2):

الإرهاق الناتج عن الاستخدام المفرط للاجتماعات المرئية مثل زوم.

Rationale: Fully descriptive translation requiring no prior knowledge; appropriate when term appears only once and immediate comprehension is prioritized over introducing new terminology.

False Friends – “Toxic” (31%):

Source (Netflix synopsis): “A bold high school student challenges the school’s toxic social hierarchy.”

Student translation:

طالبة ثانوية جريئة تتحدى التسلسل الهرمي الاجتماعي السام في المدرسة.

Problem: “Toxic” translated as “سام” (poisonous) which in Arabic retains primarily literal meaning related to actual poisons. In contemporary English, “toxic” has extended

metaphorically to describe harmful, manipulative, emotionally destructive behavior patterns. Literal “سام” applied to social hierarchy sounds confusing, suggesting chemical contamination rather than psychological harm.

Suggested corrections: Option 1: التسلسل الهرمي الاجتماعي الضار (harmful social hierarchy); Option 2: التسلسل الهرمي الاجتماعي غير الصحي (unhealthy social hierarchy); Option 3: التسلسل الهرمي الاجتماعي السلبي (negative social hierarchy). Rationale: Each option captures the metaphorical meaning without literal toxicity confusion.

5.3 Technical/Platform-Specific Errors (62%)

Technical errors reflected insufficient understanding of platform conventions.

Character Limit Violations (34%):

Source (Twitter, 67 characters): “Quick update: Our app just hit 1 million downloads! Thank you all! 🎉”

Student translation (128 characters):

تحديث سريع: لقد وصل تطبيقنا للتو إلى مليون عملية تنزيل! شكراً جزيلاً لكم جميعاً على دعمكم المستمر! 🎉

Suggested correction (89 characters):

تحديث: تطبيقنا وصل لمليون تنزيل! شكراً لكم جميعاً! 🎉

Hashtag Mishandling (29%):

Source: #foodie #homecooking #healthyliving #recipes #instafood

Student translation: #طعام #طبخ منزلي #حياة صحية #وصفات #طعام_إنستغرام

Problems: The component-by-component translation ignores established Arabic hashtag communities; the translated hashtags may not connect to existing conversation streams.

Suggested correction:

#عاشق_الطعام #طبخ_منزلي #حياة_صحية #وصفات #تصوير_الطعام #اكل_صحي

5.4 Pragmatic Errors (44%)

Tone Mismatch (Fitness App):

Source: “Hey champ! You crushed yesterday’s workout! Ready to keep that momentum going? Let’s do this! 💪”

Student translation:

مرحباً أيها البطل! لقد أنجزت تمرين الأمس بنجاح! هل أنت مستعد لمواصلة هذا التقدم؟ لنقم بهذا! 💪

Problem: While grammatically correct MSA, the translation sounds formal and distant rather than energetic and motivational; لنقم بهذا is particularly stilted.

Suggested correction:

مرحباً أيها البطل! أداؤك في تمرين الأمس كان رائعاً! هل أنت مستعد لمواصلة هذا الحماس؟ لنبدأ! 💪

6. Discussion

6.1 Error Patterns, NMT Tendencies, and the Limits of Causal Attribution

The most striking finding is the systematic correspondence between student error patterns and documented NMT tendencies, particularly in register management and cultural adaptation. When students produced over-formalised translations such as ‘لنقم بهذا’ for the casual fitness application encouragement ‘Let’s do this!’, their choices are consistent with the output tendencies of Google Translate and similar systems documented by Sajjad et al. (2020).

However, as acknowledged in Section 4.3, individual error attribution remains probabilistic. The pattern-matching inference—that systematic over-formalisation across multiple students suggests shared AI mediation rather than independent human error—is plausible but not conclusive. Interview data support the inference as eight of ten interviewed students reported using AI tools as a primary drafting resource, and their descriptions of evaluation practice reveal a consistent gap. Future research could strengthen causal claims by comparing error rates among identifiable AI-tool users versus those who report using no AI assistance, or by administering a parallel task under monitored conditions to a subset of participants.

***Student A:** “I used Google Translate for the initial translation, then corrected some obvious errors like grammar mistakes. If the translation looked good, I left it.”*

***Student B:** “I used ChatGPT and asked it to translate the text. Sometimes I changed some words to sound better, but mostly the translation was good.”*

These representative responses reveal critical gaps in AI literacy. Students evaluated output on surface grammatical correctness rather than functional appropriateness, register match, cultural adaptation, or platform suitability. None of the interviewed students articulated awareness of NMT training data composition or systematic over-formalisation tendencies. This pattern is consistent with Kenny’s (2022) account of post-editing under-correction, where users fail to identify MT-induced errors beyond surface grammatical mistakes.

The Arabic sociolinguistic context may compound this tendency. Morocco’s diglossic environment, in which Darija (Moroccan colloquial Arabic) functions as the informal register and Modern Standard Arabic as the formal register (Ferguson, 1959; Ennaji, 2005), may reinforce students’ inclination to map digital media content, however casual, onto formal MSA structures. In addition, NMT systems trained predominantly on formal Arabic corpora make the same mapping. The result is therefore a convergent rather than coincidental pattern. It is worth noting, however, that this diglossic tendency could be leveraged pedagogically. As a suggestion, students’ Darija competence can provide intuitive access to informal register that, with appropriate instruction, could be channelled productively into register-sensitive MSA output.

6.2 Critical AI Literacy as Core Competency

These findings call for a fundamental reorientation of translation pedagogy. Rather than treating AI tools as peripheral or prohibited, curricula must develop critical AI literacy as a core competency structured around five areas of understanding:

- **NMT training data composition.** Students must understand why formal Arabic corpora dominate commercial NMT training data and how this creates systematic over-formalisation tendencies in informal text domains.

- **Register mismatch recognition.** Students require explicit criteria for identifying when grammatically correct AI output is functionally inappropriate due to register, tone, or formality mismatch.
- **Cultural adaptation limitations.** Students must understand that AI systems tend to transfer cultural references literally rather than adapting them, and that communicative equivalence is a human judgement call, not an algorithmic output.
- **Platform convention awareness.** Students must know that AI systems are generally unaware of character limits, hashtag community conventions, interface terminology standards, and navigation patterns.
- **Differential tool capabilities.** Students should understand that LLMs such as ChatGPT and GPT-4 perform differently from statistical NMT systems on register-sensitive tasks (Hendy et al., 2023; Jiao et al., 2023), and that tool selection is itself a professional decision requiring informed judgement.

Strategic AI engagement requires pedagogy to move beyond the binary of prohibition versus unrestricted use, developing instead a structured workflow that includes pre-translation analysis of source text register and platform constraints, critical output evaluation against explicit criteria beyond grammatical accuracy, targeted post-editing of predictable AI tendencies, and final quality assurance against communicative purpose.

7. Pedagogical Implications

The findings suggest five interconnected curricular interventions for digital media translation programmes.

Curriculum integration. Translation competencies should be integrated systematically across the curriculum rather than siloed in a single course. A sequenced approach might introduce basic translation and cultural adaptation in Year 1, dedicate a Year 2 module to digital content localisation (cultural adaptation frameworks, platform-specific conventions, multimodal translation), and culminate in Year 3–4 capstone projects replicating professional digital media localisation workflows with client briefs, revision rounds, and AI tool integration (Kiraly, 2015).

Explicit AI literacy instruction. Given the ubiquity of AI tools in digital media workflows, curricula must address NMT capabilities and limitations directly (Kenny, 2022). This includes: foundational AI literacy covering how NMT systems work; quality evaluation frameworks providing systematic criteria for assessing AI output on register, cultural adaptation, terminology, and platform suitability; comparative analysis exercises requiring students to submit identical source texts to multiple systems (Google Translate, ChatGPT) and analyse the outputs; and post-editing practice developing efficient correction workflows (Koponen, 2016). Standard reference resources such as the Microsoft Arabic Language Style Guide, Apple Arabic Human Interface Guidelines, and major platform glossaries should be integrated into practical exercises.

Register awareness training. Explicit instruction must address the MSA register continuum from formal (legal documents, classical literature) to informal (contemporary social media). Contrastive analysis of formal and informal MSA translations of identical source texts, supported by examination of authentic Arabic digital media content—social media influencers, YouTube channels and mobile application interfaces—develops the register sensitivity students currently lack. A structured register decision-making framework should require students to justify register choices against translation skopos or purpose, target

audience, and platform conventions rather than defaulting to formal MSA. Research on L1-mediated comprehension strategies suggests that students' Darija competence may be productively engaged during the reading and comprehension phase, with MSA reserved for the translation output stage (Machaal, 2023).

Cultural adaptation decision-making. A structured decision framework adapted from Dickins et al. (2017) and Nord (2018) can move students beyond intuitive choices toward professionally defensible strategies. The framework asks students to: (1) identify the cultural element requiring a decision; (2) analyse its function (decorative, essential to meaning, humorous, brand identity marker); (3) assess target audience familiarity; (4) determine the skopos (documentary or instrumental translation); (5) select a strategy with explicit justification (preservation with explanation, functional equivalent, cultural substitution, compensatory omission); and (6) execute the translation.

Platform and terminology management. Systematic instruction should connect students' existing digital media knowledge with platform-specific translation practice. Character limit analysis, interface terminology standardisation using authoritative industry sources (Microsoft Language Portal, 2024), and parallel text research examining how professional localisers handle conventional technology terminology should be core exercises. Glossary creation for all translation projects develops the disciplined documentation habits essential for professional practice (Bowker, 2015).

8. Limitations and Future Research

Several limitations bound the generalisability of this study's findings. The single-institution, single-cohort design limits extrapolation to translation programmes at other Moroccan universities, other Arab countries, or different educational systems. Without process data from keystroke logging or screen recording (Jakobsen & Jensen, 2008), individual error attribution to human versus AI sources remains probabilistic, as explicitly acknowledged throughout. The study examined only English-to-Arabic translation; Arabic-to-English translation or other language pairs may produce different error patterns and competency gaps. Morocco's distinctive diglossic and triglossic sociolinguistic context (Ferguson, 1959; Ennaji, 2005; Youssi, 1995) may have influenced error patterns in ways not fully separable from NMT effects.

Future research priorities include: (1) controlled process studies using keystroke logging and screen recording to isolate AI tool usage and its relationship to error patterns; (2) longitudinal studies tracking students from initial training through professional practice to understand how competencies develop and how educational interventions affect long-term performance; (3) reception studies investigating how Arabic audiences across different regions respond to various translation approaches for digital media content; (4) comparative studies examining how Moroccan versus other Arab students handle digital media translation under controlled conditions; and (5) pedagogical intervention effectiveness studies using pre-post experimental designs to validate the proposed curricular changes.

9. Conclusion

Moroccan translation students bring solid Modern Standard Arabic foundations and digital media familiarity to their translation training. Yet cultural adaptation errors (73%), lexical errors (68%), and technical and platform-specific errors (62%) reveal systematically insufficient higher-order competencies: cultural mediation, strategic decision-making, and the integration of digital media knowledge with translation practice.

Error patterns are consistent with documented NMT tendencies, particularly the over-formalisation bias that stems from Arabic NMT training data composition. Students' acceptance of grammatically accurate but functionally inappropriate AI output reveals a critical gap in technological literacy. Morocco's diglossic context, in which MSA is the formal register and Darija the informal variety, may compound this tendency. Crucially, however, it also provides a pedagogical resource: students' intuitive sensitivity to informal register in Darija can, with appropriate instruction, inform register-sensitive MSA production.

Addressing these challenges requires translation programmes to: integrate translation competencies systematically across curricula; develop explicit frameworks connecting platform knowledge with translation practice; teach AI literacy and critical evaluation as core competencies, with specific attention to why Arabic NMT systems produce systematically formal output in informal text domains; explicitly address the MSA register continuum; provide cultural adaptation frameworks moving beyond binary choices; implement systematic terminology management training; and adopt authentic assessments simulating professional digital media workflows.

As AI translation technologies continue to improve, with LLMs already narrowing but not eliminating the register gap documented in earlier NMT systems, the professional translator's distinctive value lies increasingly in strategic, cultural, and pragmatic judgement that machines cannot reliably provide. Translation pedagogy must prepare students for this reality by developing the human competencies that complement rather than compete with AI capabilities.

References

- Bouamor, H., Sajjad, H., Abdelali, A., & Durrani, N. (2018). The MADAR Arabic dialect corpus and lexicon. In *Proceedings of the 11th Language Resources and Evaluation Conference* (pp. 3387–3396). European Language Resources Association. <https://doi.org/10.6317/4wv9mx9ftkrk>
- Bowker, L. (2015). Computer-aided translation: Translator training. In S.-W. Chan (Ed.), *The Routledge encyclopedia of translation technology* (pp. 88–104). Routledge.
- Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://doi.org/10.1191/1478088706qp063oa>
- Castilho, S., Doherty, S., Gaspari, F., & Moorkens, J. (2022). Approaches to human and machine translation quality assessment. In F. Gaspari, J. Moorkens, S. Castilho, S. Doherty, & Y. Erlandsson (Eds.), *Translation quality assessment: From principles to practice* (pp. 9–38). Springer. https://doi.org/10.1007/978-3-319-91241-7_2
- Darwish, K., Mubarak, H., & Abdelali, A. (2021). Arabic dialect identification. In N. Habash, S. Bouamor, & H. Hajj (Eds.), *Arabic natural language processing* (pp. 31–64). Springer. <https://aclanthology.org/2021.wanlp-1.1/>
- Dickins, J., Herve, S., & Higgins, I. (2017). *Thinking Arabic translation: A course in translation method* (3rd ed.). Routledge. <https://doi.org/10.4324/9781315471570>
- Englund Dimitrova, B. (2005). Expertise and explicitation in the translation process. John Benjamins. <https://doi.org/10.1075/btl.64>
- Ennaji, M. (2005). *Multilingualism, cultural identity and education in Morocco*. Springer. <https://doi.org/10.1007/b104063>

- Faiq, S. (2004). The cultural encounter in translating from Arabic. In S. Faiq (Ed.), *Cultural encounters in translation from Arabic* (pp. 1–13). Multilingual Matters. <https://doi.org/10.2307/jj.29308429.5>
- Farghal, M., & Shunnaq, A. (2011). *Translation with reference to English and Arabic: A practical guide* (2nd ed.). Irbid National University.
- Ferguson, C. A. (1959). Diglossia. *Word*, 15(2), 325–340. <https://doi.org/10.1080/00437956.1959.11659702>
- Göpferich, S. (2009). Towards a model of translation competence and its acquisition: the longitudinal study *Copenhagen studies in language*. 2009, Num 37, pp 11-37, 27 p
- Göpferich, S., Jakobsen, A. L., & Mees, I. M. (Eds.). (2008). Looking at eyes: Eye-tracking studies of reading and translation processing. *Samfundslitteratur*. [Looking at Eyes: Eye-Tracking Studies of Reading and Translation Processing - CBS Research Portal](https://doi.org/10.1075/btl.127)
- Hansen, G. (2010). Translation ‘errors’. In Y. Gambier & L. van Doorslaer (Eds.), *Handbook of translation studies* (Vol. 1, pp. 385–388). John Benjamins. <https://doi.org/10.1075/btl.1.tra3>
- Hendy, A., Abdelrehim, M., Sharaf, A., Raunak, V., Gabr, M., Matsushita, H., Kim, Y. J., Afli, H., & Awasthi, A. (2023). How good are GPT models at machine translation? A comprehensive evaluation. *arXiv*. <https://doi.org/10.48550/arXiv.2302.09210>
- Hurtado Albir, A. (2017). Researching translation competence. In C. V. Angelelli & B. J. Baer (Eds.), *Researching translation and interpreting* (pp. 49–58). Routledge. <https://doi.org/10.1075/btl.127>
- Husni, R., & Newman, D. L. (2015). *Arabic–English–Arabic translation: Issues and strategies*. Routledge. <https://doi.org/10.4324/9781315773339>
- Jakobsen, A. L., & Jensen, K. T. H. (2008). Eye movement behaviour across four different types of reading task. In S. Göpferich, A. L. Jakobsen, & I. M. Mees (Eds.), *Looking at eyes: Eye-tracking studies of reading and translation processing* (pp. 103–124). *Samfundslitteratur*. [Eye Movement Behaviour Across Four Different Types of Reading Task - CBS Research Portal](https://doi.org/10.1075/btl.127)
- James, C. (2013). *Errors in language learning and use: Exploring error analysis*. Routledge. <https://doi.org/10.4324/9781315842912>
- Jiao, W., Wang, W., Huang, J.-T., Wang, X., & Tu, Z. (2023). Is ChatGPT a good translator? Yes with GPT-4 as the engine. *arXiv*. <https://doi.org/10.48550/arXiv.2301.08745>
- Jiménez-Crespo, M. A. (2013). *Translation and web localisation*. Routledge. <https://doi.org/10.4324/9780203067079>
- Kenny, D. (2022). *Machine translation for everyone: Empowering users in the age of artificial intelligence*. Translation and Multilingual Natural Language Processing 18. Berlin: Language Science Press. <https://doi.org/10.17234/HIERONYMUS.9.5>
- Kiraly, D. (2015). Occasioning translator competence: Moving beyond social constructivism toward a postmodern alternative to instructionism. *Translation and Interpreting Studies*, 10(1), 8–32. <https://doi.org/10.1075/tis.10.1.02kir>
- Koehn, P., & Knowles, R. (2017). Six challenges for neural machine translation. In *Proceedings of the First Workshop on Neural Machine Translation* (pp. 28–39). Association for Computational Linguistics. <https://doi.org/10.18653/v1/W17-3204>

- Koponen, M. (2016). Machine translation post-editing and effort: Empirical studies on the post-editing process [Doctoral dissertation, University of Helsinki]. Helsinki University Digital Repository. <http://urn.fi/URN:ISBN:978-951-51-1975-9>
- Läubli, S., Fishel, M., Massey, G., Ehrensberger-Dow, M., & Volk, M. (2020). Assessing the usability of raw machine translated output. *Machine Translation*, 34(1), 1–21. <https://aclanthology.org/2013.mtsummit-wptp.10/>
- Lommel, A., Uszkoreit, H., & Burchardt, A. (2014). Multidimensional quality metrics (MQM): A framework for declaring and describing translation quality metrics. *Tradumàtica*, 12, 455–463. <https://doi.org/10.5565/rev/tradumatica.77>
- Machaal, B. (2023). The mediating role of first language during reading comprehension for translation from third language into second language: A practitioner inquiry. In *Proceedings of the 7th International Conference on New Approaches in Education* (pp. 88–96). Diamond Scientific Publishing. <https://doi.org/10.33422/7th.icnaeducation.2023.09.001>
- Massey, G., & Ehrensberger-Dow, M. (2017). Machine learning: Implications for translator education. *Lebende Sprachen*, 62(2), 300–312. <https://doi.org/10.1515/les-2017-0021>
- Mellinger, C. D., & Hanson, T. A. (2020). *Quantitative research methods in translation and interpreting studies*. Routledge. <https://doi.org/10.4324/9781003127970-27>
- Microsoft Language Portal. (2024). Arabic terminology and language style guide. Microsoft Corporation. <https://www.microsoft.com/en-us/language>
- Nagoudi, E. M. B., Elmadany, A., Abdul-Mageed, M., Alhindi, T., & Cavusoglu, H. (2022). Machine translation from Arabic dialects: A closer look at dialectal Arabic. *arXiv*.
- Nord, C. (2018). *Translating as a purposeful activity: Functionalist approaches explained* (2nd ed.). Routledge.
- PACTE. (2003). Building a translation competence model. In F. Alves (Ed.), *Triangulating translation: Perspectives in process-oriented research* (pp. 43–66). John Benjamins. <https://doi.org/10.1075/btl.45.06pac>
- PACTE. (2005). Investigating translation competence: Conceptual and methodological issues. *Meta*, 50(2), 609–619. <https://doi.org/10.7202/011004ar>
- Saadane, H., & Habash, N. (2015). A conventional orthography for Algerian Arabic. In *Proceedings of the Second Workshop on Arabic Natural Language Processing* (pp. 69–79). Association for Computational Linguistics. <https://doi.org/10.18653/v1/W15-3208>
- Sajjad, H., Abdelali, A., Durrani, N., & Dalvi, F. (2020). AraBench: Benchmarking dialectal Arabic–English machine translation. In *Proceedings of the 28th International Conference on Computational Linguistics (COLING 2020)* (pp. 5094–5107). International Committee on Computational Linguistics. <https://aclanthology.org/2020.coling-main.447>
- Vilar, D., Xu, J., D’Haro, L. F., & Ney, H. (2006). Error analysis of statistical machine translation output. In *Proceedings of the Fifth International Conference on Language Resources and Evaluation* (pp. 697–702). European Language Resources Association. <https://doi.org/10.63317/453fcbexxngk>
- Youssi, A. (1995). The Moroccan triglossia: Facts and implications. *International Journal of the Sociology of Language*, 112(1), 29–43. <https://doi.org/10.1515/ijsl.1995.112.29>

Appendix A: Source Texts

Text 1 — Twitter/X (Technology Company Product Announcement)

“Big news! We’re rolling out dark mode across all platforms 🌙” / “We heard you! Dark mode is finally here.” / “Your eyes will thank you during those late-night scrolling sessions.”

Text 2 — Instagram Caption (Fashion Influencer)

“Giving major cottage core vibes with this thrifted find! 🌸 Sustainable fashion is the new black, and I’m here for it ❤️”

Text 3 — YouTube Description (Cooking Channel)

“Join me every Tuesday for Taco Tuesday! We’ll explore authentic Mexican street food recipes that’ll transport you straight to Guadalajara.”

Text 4 — Netflix Synopsis (Teen Drama)

“When her perfect life crumbles, a privileged teen must navigate public school for the first time. As she confronts a toxic social hierarchy, she discovers unexpected friendships.”

Text 5 — Fitness App Interface

Push notification: “Hey champ! You crushed yesterday’s workout! Ready to keep that momentum going? Let’s do this! 💪” Settings menu: “Push Notifications | Reminder Time | Language”

Text 6 — Facebook Post (Environmental NGO)

“Every small action counts. 🌍 By reducing plastic use, we can make a massive difference together. Join our #PlasticFreeChallenge this month.”

Text 7 — TikTok Spoken Transcript

“POV: You’re a millennial trying to explain to Gen Z why we still use the laughing-crying emoji 😂 Not me having a whole existential crisis over this!”

Text 8 — Remote Work Blog Excerpt

“WFH has its perks — no commute, comfy clothes, and more flexibility. But it also means dealing with Zoom fatigue and blurred work-life boundaries.”

Appendix B: Error Distribution by Genre

Table 2. Error Patterns Across Digital Media Genres (% of translations in genre showing error category)

Genre	Cultural Adapt.	Lexical	Technical	Syntactic	Pragmatic	Omission/ Addition	Orthographic
Twitter/X	61	58	71	44	38	29	22
Instagram	89	74	55	47	51	36	31
YouTube	76	63	48	52	42	38	24
Netflix	71	69	44	58	47	33	27
Fitness App	58	72	82	49	64	31	29
Facebook NGO	67	61	53	48	36	28	22
TikTok	84	76	68	54	58	42	36
Blog Post	59	67	49	56	41	31	29
Overall	73	68	62	51	44	33	28

Note. Genre-specific percentages reflect the proportion of translations within each genre (n = 45 per genre) showing at least one error in the given category. Overall row reproduces Table 1 totals.