

Assessing Student Translator Competence in Digital Tourism Texts: A Rubric-Based Comparison with Machine Translation

Mohammed Hussein A. Alahmadi

College of Languages and Translation, Taibah University, Saudi Arabia
mhaahmadi@taibahu.edu.sa

Abstract

This mixed-methods study examines Arabic translations of a digital tourism text as a skopos-driven task and compares senior students' performance with neural machine translation (NMT). Thirty students (15 male and 15 female) translated the approximately 300-word promotional text "AlUla Calling." Two trained raters assessed the translations with a five-dimension analytic rubric covering idiomaticity, cultural and stylistic transfer, digital discourse adaptation, semantic accuracy, and fluency/readability. Students achieved a mean total of 16.60/ 20 (SD = 3.22, 95% CI [15.39, 17.80]). Their strongest dimensions were idiomaticity (M = 3.53) and digital discourse adaptation (M = 3.47); fluency/readability was weakest and most variable (M = 2.87, SD = 1.00). Google Translate and DeepL outputs, used as task-specific benchmarks, scored 18.00/20. NMT was stronger in semantic accuracy, fluency, and digital discourse adaptation, whereas students scored higher in idiomaticity and cultural-stylistic transfer. Thematic analysis identified literal phrasing, register drift, inconsistent treatment of digital features, and minor semantic errors, alongside effective idiomatic substitution, localization, creative adaptation, and compensation. The findings support instruction that combines fluency and meaning-control practice with critical NMT post-editing while preserving the culturally resonant voice that human translators contribute.

Keywords: Digital Tourism, Skopos Theory, Translation Pedagogy, Student Translators, Rubric-Based Assessment, Saudi Vision 2030, Machine Translation.

1. Introduction

Saudi Vision 2030 has placed tourism at the center of national economic diversification, creating sustained demand for multilingual communication about destinations such as AlUla, NEOM, the Red Sea, and Diriyah (Saudi Vision 2030, 2024). Much of that communication now occurs through websites, social media, mobile applications, and user-generated platforms. These channels do more than convey information: they promote destinations through direct address, informal register, hashtags, emoji, sensory imagery, and calls to action. Digital tourism discourse therefore combines the persuasive function of traditional tourism language with the conversational and participatory conventions of online communication (Dann, 1996; Francesconi, 2014; Manca, 2016; Thurlow & Jaworski, 2010; Zappavigna, 2018).

This hybrid genre complicates translation. Idioms and colloquialisms must sound natural without

losing their persuasive force; cultural and popular-culture references may require explanation or adaptation; and platform features such as hashtags, emoji, abbreviations, and interactive questions must retain their pragmatic role. Translators must also balance informality with credibility and preserve accurate information about places and institutions. Effective work consequently draws on bilingual competence, cultural mediation, digital literacy, and transcreation rather than formal correspondence alone.

These demands are well explained by functional and competence-based approaches. Skopos theory treats translation as purposeful action and evaluates a target text according to its intended function and audience (Nord, 1997; Reiss & Vermeer, 1984/2014; Vermeer, 1989/2012). Because tourism promotion should read as audience-facing communication rather than an overtly translated document, it calls for the covert orientation described by House (1997). The PACTE model likewise presents translation competence as an interaction among bilingual, extra-linguistic, instrumental, knowledge-about-translation, and strategic sub-competences (Hurtado Albir, 2015, 2017; PACTE Group, 2003, 2009, 2015). Kelly (2005) emphasizes cultural mediation and professional practice, while González-Davies (2014) connects competence development to authentic tasks, reflection, and collaborative problem-solving.

Translation quality assessment should therefore capture functional success as well as error frequency. Holistic and analytic rubrics can make multiple quality dimensions visible and provide diagnostic feedback linked to teachable skills (Waddington, 2001). In the present study, idiomaticity and fluency indicate target-language control; cultural-stylistic transfer reflects intercultural and extra-linguistic competence; digital discourse adaptation reflects genre and platform awareness; and semantic accuracy reflects controlled meaning transfer. Together, the dimensions operationalize the task's skopos: persuasive, natural, and culturally accessible Arabic tourism communication.

NMT provides an important comparison. Transformer architectures have improved grammaticality and semantic adequacy across many domains (Vaswani et al., 2017), but research continues to identify problems when functional success depends on document-level coherence or creative, rhetorically sensitive choices (Läubli et al., 2018; Taivalkoski-Shilov & Koponen, 2023; Toral & Way, 2018). Digital tourism is therefore a useful test of complementary strengths: machines may offer fluent and stable wording, while humans may better interpret communicative intent and construct culturally resonant voice.

The literature also shows why digital features cannot be treated as decorative extras. Direct questions and imperatives position readers as participants, while hashtags index content and invite circulation (Thurlow & Jaworski, 2010; Zappavigna, 2018). Emoji can likewise convey stance, playfulness, or emphasis in very little space (Padilla, 2023). Their appropriate translation depends on the platform, the surrounding visuals, and target-audience expectations. A formally accurate sentence may therefore be functionally weak if it removes the invitation to interact or replaces a peer-like voice with institutional prose. Conversely, creative adaptation remains accountable to semantic constraints: promotional energy cannot justify altered place names, inaccurate institutional references, or

misleading claims.

This balance has implications for translator education. Digital promotional work is not fully represented by curricula centered on literary, legal, or technical genres, where creative latitude and platform interaction differ. Authentic, task-based projects can require students to define audience and purpose, explain strategic choices, and test whether a target text works alongside visual and platform constraints (González-Davies, 2014). Peer review and formative rubrics make competing criteria explicit, while technology integration reflects professional workflows in which translators must evaluate rather than merely accept MT output (Kelly, 2005; Waddington, 2001). These practices develop the strategic control that PACTE identifies as coordinating accuracy, cultural knowledge, resource use, and problem -solving.

Empirical work on student performance in this genre remains limited, particularly in Saudi translator education. This study addresses that gap by assessing senior students with a specialized functional rubric, comparing their performance with Google Translate and DeepL benchmarks, and identifying competence gaps that can inform pedagogy in a tourism sector shaped by Vision 2030.

2. Research Questions

The study asks: (RQ1) To what extent do senior translation students demonstrate competence in idiomaticity, cultural and stylistic transfer, digital discourse adaptation, semantic accuracy, and fluency/readability? (RQ2) How does their performance compare with Google Translate and DeepL across these dimensions and overall quality? (RQ3) Which recurrent errors and strategy choices reveal competence gaps, and what pedagogical implications follow?

3. Methods

3.1 Design and Participants:

A descriptive comparative mixed-methods design combined rubric scores with thematic analysis. Participants were 30 senior undergraduates in the BA Translation program at Taibah University, Yanbu Branch, during fall 2025. The cohort included 15 men and 15 women, all native Arabic speakers studying English as an L2; their ages ranged from 21 to 24 years ($M = 22.3$, $SD = 0.9$). Their translation experience was primarily academic. Participation was voluntary, withdrawal was permitted without penalty, and work was reported anonymously.

Two Arabic-speaking applied linguists independently assessed all translations. One had more than 10 years of teaching experience and the other approximately 8 years; both had experience in translation teaching and rubric-based assessment. A two-hour training session reviewed the translation brief, rubric descriptors, and scoring procedures. Disagreements were resolved by discussing textual evidence against the descriptors, with researcher adjudication when necessary.

3.2 Task, Rubric, and NMT Benchmarks:

Students translated the approximately 300-word English text "AlUla Calling" (Appendix A),

designed to resemble a social-media tourism post. It included direct address, rhetorical questions, sensory imagery, idioms such as "game-changer," the abbreviation "TBH," the cultural reference "Indiana Jones," and tourism terminology such as "UNESCO World Heritage Site." Under supervised exam conditions, students had 90 minutes to produce persuasive Arabic digital copy. Dictionaries were allowed, but machine translation and online paraphrasing were prohibited.

The analytic rubric (Appendix B) was adapted from function-oriented translation assessment and research on tourism and digital discourse (Agorni, 2012; Manca, 2016; Thurlow & Jaworski, 2010; Waddington, 2001; Zappavigna, 2018). Five criteria -idiomaticity, cultural and stylistic transfer, digital discourse adaptation, semantic accuracy, and fluency/readability--were each scored from 1 (very weak) to 4 (excellent), for a maximum of 20. Two experts in translation and digital tourism reviewed the rubric for relevance and clarity, and their suggested revisions were incorporated. A separate pilot was not feasible because the available advanced cohort comprised the study sample.

Each criterion represented a distinct but overlapping aspect of functional competence. Idiomaticity captured natural Arabic solutions to colloquial and figurative language; cultural-stylistic transfer assessed references, stance, and promotional voice; and digital discourse adaptation examined direct address, interaction, and platform markers. Semantic accuracy addressed information, nuance, and unintended distortion, while fluency/readability covered grammar, cohesion, and the extent to which the result read as publishable Arabic copy. Using separate dimensions allowed a translation to receive credit for persuasive adaptation without concealing errors in meaning or language control.

For comparison, the complete source text was entered into the default web interfaces of Google Translate and DeepL on 15 October 2025. Outputs were neither corrected nor post-edited and were scored with the same rubric. Because each system supplied only one output for one source text, their scores are illustrative task-specific benchmarks, not samples supporting population-level inference.

3.3 Analysis:

Descriptive statistics summarized scores for the 30 student translations. The same rubric provided a dimension-by-dimension comparison with the two NMT outputs. Qualitative thematic coding, guided by the rubric and Patton (2015), identified recurring errors and strategies affecting the persuasive skopos. Themes are reported as functional tendencies rather than exhaustive frequencies.

4. Results

4.1 Quantitative Results:

Students achieved generally strong scores (Table 1). Idiomaticity was highest ($M = 3.53$, $SD = 0.63$), followed by digital discourse adaptation ($M = 3.47$, $SD = 0.63$), cultural-stylistic transfer ($M = 3.40$, $SD = 0.72$), and semantic accuracy ($M = 3.33$, $SD = 0.71$). Fluency/readability was lowest and most variable ($M = 2.87$, $SD = 1.01$). The mean total was 16.60/20 ($SD = 3.22$, range = 9-20, 95% CI [15.39, 17.80]).

Table (1): Student translator descriptive scores

Dimension	N	Mean	SD	Median	Min	Max	95% CI
Idiomacity	30	3.533	0.629	4.0	2	4	[3.299, 3.768]
Cultural & stylistic	30	3.400	0.724	4.0	2	4	[3.130, 3.670]
Digital discourse	30	3.467	0.629	4.0	2	4	[3.232, 3.701]
Semantic accuracy	30	3.333	0.711	3.0	2	4	[3.068, 3.599]
Fluency & readability	30	2.867	1.008	2.5	1	4	[2.490, 3.243]
Total (/20)	30	16.600	3.223	16.0	9	20	[15.397, 17.803]

4.2 Student-NMT Comparison:

The NMT benchmark total (18.00/20) exceeded the student mean (16.60/20), but the dimensions reveal different profiles (Table 2). Students scored higher in idiomacity (3.53 vs. 3.00) and cultural-stylistic transfer (3.40 vs. 3.00). NMT scored higher in digital discourse adaptation (4.00 vs. 3.47), semantic accuracy (4.00 vs. 3.33), and fluency/readability (4.00 vs. 2.87). The comparison therefore suggests stronger local voice and cultural engagement among students but greater accuracy and surface fluency in the two machine outputs.

Table (2): Student and NMT mean scores

Dimension	Student mean	NMT mean
Idiomacity	3.53	3.00
Cultural & stylistic	3.40	3.00
Digital discourse	3.47	4.00
Semantic accuracy	3.33	4.00
Fluency & readability	2.87	4.00
Total (/20)	16.60	18.00

4.3 Qualitative Patterns:

Literal transfer was a recurrent source of awkwardness. Some students recreated expressions such as "more than just a pretty picture" naturally, whereas others rendered idioms, "travel feed," or social-media collocations compositionally, preserving recoverable meaning but reducing idiomacity and readability. Cultural references were usually retained, yet names such as Indiana Jones and UNESCO were often merely transcribed. Brief contextualization sometimes improved accessibility, while inconsistent renderings of AlUla or Hegra weakened accuracy and destination branding.

Register also varied. Some translations adopted a generic, formal brochure style and lost the source's conversational immediacy; others mixed colloquial and formal Arabic unevenly. Stronger texts preserved direct address, rhetorical questions, calls to action, emoji, and locally meaningful versions of the closing hashtag. Weaker texts omitted or neutralized these features, reducing the post's platform-native feel. Semantic problems were usually local rather than wholesale: errors involved proper names, sequence, intensity, or unintended connotations, including a rendering equivalent to "rest in peace" for relaxing beneath the stars.

The analysis also identified productive strategies. Higher-scoring students used idiomatic

substitution and Saudi or Gulf discourse markers to recreate the energetic voice, while safer paraphrases conveyed meaning with less rhetorical impact. Cultural adaptation occurred more often through audience-facing voice and brief explication than through replacement of globally recognizable references. Creative hashtag and emoji choices preserved pragmatic function. When a particular idiom could not be reproduced, some students compensated elsewhere through stronger calls to action, sensory detail, or emotional framing. These choices show purposeful attention to the overall skopos despite uneven linguistic control.

The NMT outputs were grammatical, meaning-faithful, and consistent with digital conventions, matching their ceiling scores for semantic accuracy, fluency/readability, and digital discourse adaptation. Their phrasing was nevertheless safer and more standardized, explaining the lower idiomaticity and cultural-stylistic ratings. This complementary pattern must remain task-specific because only one text and one output per system were assessed.

5. Discussion

The results indicate a meaningful split between persuasive voice and linguistic control. Students generally recognized that the text should function as engaging Arabic tourism content rather than a literal record of the English source. Their high idiomaticity and digital-discourse scores support a skopos-oriented interpretation: many could position the reader, recreate conversational energy, and make culturally informed choices. However, variability in grammar, cohesion, collocation, and semantic detail prevented some texts from reaching publishable quality. In digital promotion, fluency is not cosmetic; interruptions to the expected "scrollable" flow can reduce credibility and persuasive momentum.

The comparison with NMT sharpens this interpretation. Google Translate and DeepL produced accurate and polished copy, but their more conservative phrasing was less locally resonant. Human work was often more rhetorically ambitious, although that ambition sometimes introduced error. The relevant pedagogical lesson is not that one mode is categorically superior. Rather, machine output can supply grammatical and propositional stability, while trained translators add audience knowledge, cultural judgment, and a voice calibrated to communicative purpose.

The qualitative evidence also cautions against reducing competence to the total score. A literal rendering could remain understandable yet fail to sound like digital promotion, while an energetic adaptation could meet the persuasive purpose but introduce a small factual or grammatical defect. The analytic profile makes these trade-offs visible. Students' use of idiomatic substitution, partial explication, localized address, and compensation demonstrates strategic awareness, even when execution was inconsistent. Their errors likewise indicate where monitoring broke down: proper names were not always verified, register shifted within a text, or creative language was not followed by a final accuracy and fluency review. These are developmental problems that targeted revision can address.

Instruction should first treat fluency as editing for publication, with attention to sentence-level

accuracy, cohesion, and consistent register rather than isolated grammar drills. Second, semantic accuracy should be framed as brand-critical precision: proper nouns, institutional names, place labels, sequence, and intensity contribute directly to destination identity and trust. Verification routines and terminology checks can improve accuracy without suppressing creativity. Third, students should practice skopos-driven NMT post-editing. They can diagnose a machine draft, retain useful fluency, and then revise idioms, cultural framing, direct address, calls to action, and hashtags for a defined Arabic audience. Authentic tasks, peer review, reflective justification, and diagnostic rubric feedback can help students coordinate these decisions (González-Davies, 2014; Kelly, 2005; PACTE Group, 2015; Waddington, 2001).

For Saudi translator education, these priorities respond directly to expanding tourism and digital-marketing needs. Curricula should create sustained opportunities to work with platform-specific discourse and multimodal briefs, evaluate machine output critically, and defend translation choices in relation to audience and purpose. The goal is a professional profile that combines technological competence with the cultural and rhetorical value that the student results already demonstrate.

6. Limitations and Future Research

The study assessed one source text and only one output from each NMT system, so the machine scores are benchmarks and the findings should not be generalized across genres, systems, or prompts. A separate rubric pilot was also precluded by cohort size. Future research should use multiple tourism texts, report inter-rater reliability, include larger and more varied student samples, and compare additional engines or controlled prompting conditions. Process data could further show how students use resources, monitor meaning, and revise toward the skopos.

Longitudinal designs would be especially useful for testing whether targeted instruction changes the competence profile identified here. Researchers could compare initial translation, unassisted revision, and NMT-supported post-editing, then examine not only final rubric scores but also the reasons students give for retaining or rejecting machine suggestions. Studies using authentic multimodal posts could additionally test how images, character limits, audience analytics, and platform conventions influence translation decisions. Such evidence would clarify whether gains transfer beyond a supervised, text-only task to professional digital-tourism workflows.

7. Conclusion

Senior students already show valuable skopos-aware strengths in idiomatic and culturally engaging Arabic tourism writing, but they need stronger fluency and meaning control to produce consistently publishable digital copy. The two NMT outputs offered a useful baseline for accuracy and polish, while students contributed more locally resonant persuasion. These results frame competence as the ability to coordinate accuracy, natural language, culture, platform conventions, and communicative purpose rather than maximize any one dimension in isolation. Translator education should build on this complementarity through authentic digital tasks, targeted revision, and critical post-editing that preserves human cultural judgment. In the Saudi tourism context, that combined capacity can help

future translators produce reliable content that also sounds purposeful, inviting, and locally meaningful.

References

- Agorni, M. (2012). Questions of mediation in the translation of tourist texts. *Altre Modernità*, (20), 1–11. <https://doi.org/10.13130/2035-7680/1963>.
- Dann, G. M. S. (1996). *The language of tourism: A sociolinguistic perspective*. CABI.
- Francesconi, S. (2014). *Reading tourism texts: A multimodal analysis*. Channel View Publications. <https://doi.org/10.21832/9781845414283>.
- González-Davies, M. (2014). Towards a plurilingual development paradigm: From spontaneous to informed use of translation in additional language learning. *The Interpreter and Translator Trainer*, 8(1), 8–31. <https://doi.org/10.1080/1750399X.2014.908555>.
- House, J. (1997). *Translation quality assessment: A model revisited*. Gunter Narr Verlag.
- Hurtado Albir, A. (2015). The acquisition of translation competence. *Meta*, 60(2), 256–280. <https://doi.org/10.7202/1032857ar>.
- Hurtado Albir, A. (Ed.). (2017). *Researching translation competence by PACTE Group*. John Benjamins. <https://doi.org/10.1075/btl.127>.
- Kelly, D. (2005). *A handbook for translator trainers: A guide to reflective practice*. St. Jerome.
- Läubli, S., Sennrich, R., & Volk, M. (2018). Has machine translation achieved human parity? A case for document-level evaluation. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing* (pp. 4791–4796). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D18-1512>.
- Manca, E. (2016). *Persuasion in tourism discourse: Methodologies and models*. Cambridge Scholars Publishing.
- Nord, C. (1997). *Translating as a purposeful activity: Functionalist approaches explained*. St. Jerome.
- Padilla, X. A. (2023). Emoji code: From frequency-function interface to digital discursive identity. *Círculo de Lingüística Aplicada a la Comunicación*, 93, 243–257. <https://doi.org/10.5209/clac.83394>.
- PACTE Group. (2003). Building a translation competence model. In F. Alves (Ed.), *Triangulating translation: Perspectives in process-oriented research* (pp. 43–66). John Benjamins.
- PACTE Group. (2009). Results of the validation of the PACTE translation competence model:

- Acceptability and decision making. *Across Languages and Cultures*, 10(2), 207–230. <https://doi.org/10.1556/Acr.10.2009.2.3>.
- PACTE Group. (2015). Results of PACTE's experimental research on the acquisition of translation competence: The acquisition of declarative and procedural knowledge in translation: The dynamic translation index. *Translation Spaces*, 4(1), 29–53. <https://doi.org/10.1075/ts.4.1.02bee>.
 - Patton, M. Q. (2015). *Qualitative research & evaluation methods* (4th ed.). Sage.
 - Reiss, K., & Vermeer, H. J. (2014). *Towards a general theory of translational action: Skopos theory explained* (C. Nord, Trans.). Routledge. (Original work published 1984). <https://doi.org/10.4324/9781315759715>.
 - Saudi Vision 2030. (2024). *Vision 2030 annual report 2024*. <https://www.vision2030.gov.sa/media/r3ij40wu/en-annual-report-vision2030-2024.pdf>.
 - Taivalkoski-Shilov, K., & Koponen, M. (2023). Literary post-editing and the question of copyright. *HERMES - Journal of Language and Communication in Business*, (63), 195–207. <https://doi.org/10.7146/hjlb.vi63.137012>.
 - Thurlow, C., & Jaworski, A. (2010). *Tourism discourse: Language and global mobility*. Palgrave Macmillan.
 - Toral, A., & Way, A. (2018). What level of quality can neural machine translation attain on literary text? From principles to practice. In J. Moorkens, S. Castilho, F. Gaspari, & S. Doherty (Eds.), *Translation quality assessment* (pp. 263–287). Springer. https://doi.org/10.1007/978-3-319-91241-7_12.
 - Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł., & Polosukhin, I. (2017). Attention is all you need. In *Advances in neural information processing systems* 30 (pp. 5998–6008). Curran Associates, Inc.
 - Vermeer, H. J. (2012). Skopos and commission in translational action (A. Chesterman, Trans.). In L. Venuti (Ed.), *the translation studies reader* (3rd ed., pp. 191–202). Routledge. (Original work published 1989).
 - Waddington, C. (2001). Different methods of evaluating student translations: The question of validity. *Meta*, 46(2), 311–325. <https://doi.org/10.7202/004583ar>.
 - Zappavigna, M. (2018). *Searchable talk: Hashtags and social media metadiscourse*. Bloomsbury.

Appendix A. Source Text: "AlUla Calling"

AlUla Calling: More Than Just a Pretty Picture! ✨

Hey, globetrotters! If your travel feed is looking a little bland, I've got the perfect fix. Forget everything you thought you knew about desert destinations because AlUla is the real deal and a total game-changer. Seriously, this place is so much more than just an IG-worthy backdrop. 📸

When you first arrive, you'll want to hit the ground running. Start by exploring Hegra, Saudi Arabia's first UNESCO World Heritage site. The ancient tombs carved into the rocks are absolutely mind-blowing. It feels like you've stepped onto a movie set—TBH, it gives major Indiana Jones vibes, just without the snakes! 😊

But AlUla isn't stuck in the past. After a day of exploring, you can experience some epic desert "glamping." Imagine sleeping under a blanket of a million stars... it's a once-in-a-lifetime experience that'll make you want to disconnect for good.

For those who love to wander off the beaten path, the Old Town is a must-see. Its labyrinth of mudbrick houses tells stories from centuries ago. Just make sure your phone is charged because you won't want to miss a single shot.

So, are you ready to pack your bags? A trip to AlUla isn't just a vacation; it's an adventure that feeds the soul. Trust me, you'll come back with stories that your friends won't believe. What are you waiting for? #YallaAlUla 🐪

Appendix B. Translation Quality Assessment Rubric for Digital Tourism Texts

This rubric assesses student translations of digital texts, including educational and social-media content. The criteria are adapted from the assessment framework described in the study.

Criteria	4: Excellent	3: Good	2: Developing	1: Needs Improvement
1. Idiomatic Expressions	Idiomatic and colloquial expressions are translated with functionally equivalent and natural-sounding target language idioms. The translation is creative and contextually appropriate.	Most idioms are translated effectively, often through accurate paraphrase. The meaning is clear, though some idiomatic nuance may be lost.	Idioms are often translated literally or paraphrased inaccurately, resulting in awkward phrasing. The meaning may be partially obscured.	Idioms are consistently translated literally, creating nonsensical or confusing text. The source meaning is lost.
2. Cultural & Stylistic Transfer	Cultural references, humor, and stylistic register are skillfully adapted for the target audience. The translation maintains the engagement and tone of the original text.	Cultural references and humor are generally understood and transferred, but the adaptation may lack nuance or impact. The tone is mostly appropriate.	There are attempts to transfer cultural references, but they are often omitted or misinterpreted, causing confusion. Humor is typically lost.	Cultural references and humor are ignored or mistranslated, leading to a flat or inaccurate rendering. The stylistic register is inappropriate.
3. Digital Discourse Adaptation	Informal tone, abbreviations, and interactive elements (e.g., emojis) are transferred creatively and effectively, preserving the pragmatic function of the source text.	Interactive elements are transferred but may be slightly formalized. The informal tone is mostly maintained but with some inconsistencies.	The translation is over-formalized. Emojis and abbreviations are often omitted or replaced with lengthy descriptions, losing authenticity.	The text is completely formalized. All digital discourse features are removed, failing to capture the nature of the source text.
4. Semantic Accuracy	The meaning of the source text is conveyed with high precision and no distortions. Nuances in meaning are fully preserved.	The core meaning is conveyed accurately. There may be minor shifts in nuance, but they do not affect overall comprehension.	There are some inaccuracies or misinterpretations that lead to a partial loss of the intended meaning.	Significant errors in meaning fundamentally distort the message of the source text.
5. Overall Fluency & Readability	The target text is highly fluent, clear, and natural. It reads as if it were originally written in the target language, with no grammatical errors.	The target text is clear and readable with only minor grammatical or syntactic errors that do not impede understanding.	The text is understandable but contains noticeable grammatical errors and awkward sentence structures that affect readability.	The text is difficult to read due to frequent and significant grammatical errors and poor sentence construction.

Scoring and Comments

Total Score: _____ / 20

General Comments and Feedback for Improvement: