

Article

Critical Machine Translation Praxis in Tourism Discourse: Translating Latin American Destination Slogans Across AI-Based MT Engines

Pamela Doran ¹, and Carlos Zarzalejo ²

Citation: Doran, P. & Zarzalejo, C.
(2026). Critical Machine Translation
Praxis in Tourism Discourse:
Translating Latin American Destination
Slogans Across AI-Based MT Engines.
LABSREVIEW, 1(2): 24-36.
<https://doi.10.70469/labsreview.v3i1.47>.
Academic Editor: Jorge Martín Díez
Received: 10/5/2025
Revised: 1/21/2026
Accepted: 16/5/2026
Published: 7/14/2026



Copyright: © with the authors. This
Open Access article is distributed under
the terms and conditions of the Creative
Commons Attribution (CC BY4.0).

¹ SUNY Empire State University; pamela.doran@sunyempire.edu

² La Era Dígita; azarzalejo@gmail.com

* Correspondence: pamela.doran@sunyempire.edu

Abstract: Tourism slogans function as cultural invitations and are often among the first messages encountered by international travelers. Increasingly, these messages are mediated through free machine translation (MT) platforms such as Google Translate, DeepL, and Microsoft Translator. This pilot study examines whether these systems transmit not only semantic meaning but also the persuasive tone embedded in tourism discourse. To address this question, the study introduces Critical Machine Translation Praxis, an analytical approach that combines Baker's taxonomy of translation shifts with an affective semiotic perspective. Forty Latin American tourism slogans were translated across three MT engines and compared with human benchmark translations. Each output was evaluated for fidelity, translation-shift patterns, and persuasive effect. The findings reveal a notable divide: although 52.5% of the translations maintained high semantic fidelity, many of the outputs exhibited tonal flattening, excessive literalism, or culturally misleading cues. DeepL produced the most natural-sounding translations overall, while Google Translate and Microsoft Translator showed greater variability, particularly in the handling of idioms, affective nuance, and culturally embedded expressions. These results suggest that tonal drift and lexical imprecision can weaken the persuasive force of tourism slogans. The study argues that translation quality in tourism communication should not be reduced to accuracy alone; rather, emotional resonance, cultural nuance, and destination appeal remain central to how places are represented, interpreted, and received by global audiences.

Keywords: Affective blindness, Artificial Intelligence, intercultural communication, machine translation, tourism slogans.

1. Introduction

Tourism is a cornerstone of Latin America's economy, contributing billions of dollars annually and sustaining millions of jobs. In this highly competitive sector, slogans serve as strategic assets, condensing national identity and brand promise into a few memorable words. For many destinations, these short phrases are the first, and sometimes only, encounter international audiences have with a country's image. Increasingly, that encounter is mediated by free machine translation (MT) systems such as Google Translate, DeepL, and Microsoft Translator.

As tourists browse websites, social media posts, or promotional campaigns, these tools often provide the gateway through which slogans are understood. This reliance reaches beyond convenience. For travelers with limited language proficiency, or for those relying on screen readers, machine translation can determine how

accessible and inclusive a destination appears. When a translation flattens tone or distorts cultural meaning, the effect extends beyond weakened persuasion. It risks creating barriers to access and shaping perceptions in ways that exclude rather than welcome. At stake is not only linguistic accuracy, but also the preservation of persuasive force, emotional resonance, and accessibility that global tourism communication requires.

Machine translation systems have made substantial progress in grammatical and semantic accuracy, yet they continue to display what can be described as "affective blindness" or the inability to perceive and reproduce emotional undertones. In tourism marketing, persuasion relies not only on linguistic precision but also on tone, rhythm, and cultural resonance. Even small distortions can carry significant consequences. A slogan crafted to radiate warmth may appear flat, while an idiom intended to spark excitement may surface as awkward or confusing. For example, the Spanish expression "se me sale el corazón del pecho", which conveys overwhelming joy, frequently emerges in MT as "my heart comes out of my chest," a literal rendering that risks unsettling visitors rather than welcoming them.

This pilot study introduces Critical MT Praxis, an applied framework that brings together Baker's (1992) taxonomy of translation shifts with an affective semiotic lens. The framework evaluates not only structural distortions, level, category, structure, unit, and intra-system, but also the degree to which a slogan's emotional and cultural force is preserved in translation. By linking technical analysis with intercultural communication theory, Critical MT Praxis foregrounds both meaning and affect as core criteria for evaluation.

Our exploratory analysis of 40 Latin American tourism slogans highlights both the promise and the limits of contemporary MT. Just over half (52.5%) of slogans retained clarity and persuasive tone. Nearly half (47.5%), however, suffered meaningful losses through tone flattening, misleading literalism, or diluted cultural cues. These distortions matter: mistranslations can weaken brand identity, erode trust, and, in some cases, create accessibility risks for multilingual users or those relying on assistive technologies.

The central research question guiding this pilot is:

To what extent do free MT tools preserve the persuasive and affective qualities of Latin American tourism slogans, and what risks emerge when they do not?

Answering this question has both academic and practical value. For scholars, it provides an empirical case of how affective blindness manifests in short-form marketing discourse. For practitioners, it offers a straightforward pre-launch screening method: test slogans in free MT engines, identify distortions, and revise as needed.

By integrating semantic and affective dimensions, Critical MT Praxis makes significant contributions to translation studies, intercultural communication, and marketing ethics. It underscores that in tourism, language is never just informational; it is persuasive, affective, and inclusive, shaping how destinations are experienced and trusted by global audiences.

2. Literature Review

Tourism descriptions and slogans have long been recognized as a domain where language functions not only to describe destinations but also to persuade and shape identity (Dann, 1996). Slogans carry particular force, condensing cultural resonance and emotion into short, memorable lines. However, when these slogans circulate through machine translation (MT), much of their persuasive impact can be weakened or lost.

2.1 Baker's Shift Taxonomy

Baker's (1992) taxonomy of translation shifts, level, category, structure, unit, and intra-system. It offers a systematic way to document semantic distortions. While MT systems are increasingly accurate at the structural level, they struggle with affective nuance. This limitation can be described as "affective blindness", or the incapacity of neural systems to reproduce tone, warmth, or cultural resonance. Barthes's (1980) distinction between *studium* (the shared, literal meaning) and *punctum* (the affective spark that moves us) provides a valuable lens. MT often retains *studium* but erases *punctum*, flattening language into information devoid of affect. For example, "Colombia es pasión" becomes "Colombia is passion", a translation that is lexically correct but emotionally flat. Similarly, idioms such as "se me sale el corazón del pecho" appear as "my heart comes out of my chest," producing grotesque or confusing imagery rather than heartfelt excitement. These distortions are not trivial or small. They illustrate how affective meaning, central to persuasion, is lost through machine-rendered Literalism.

2.2 Intercultural Communicative Competence (ICC)

Byram (1997) defines intercultural communicative competence (ICC) as the ability to navigate across cultures with curiosity, empathy, and sensitivity. Human translators draw on ICC when choosing between “usted” and “tú,” or when reshaping metaphors so they resonate in another cultural context. MT systems, by contrast, predict word sequences without awareness of context. As House (2015) observes, this often yields surface-level accuracy but socially off-key phrasing. A typical case is Curaçao, where “siente el calor” is rendered as “feel the warmth” by a human translator, but MT typically outputs “feel the heat.” This can give the impression of discomfort instead of hospitality. ICC highlights why humans still outperform machines in contexts where persuasion depends on affect and cultural nuance.

2.3 Research Gap and Contribution

Most MT research focuses on extended texts such as literature, speeches, or idioms (Bowker & Buitrago-Ciro, 2015; Klimova, 2025). Short-form marketing discourse, despite its commercial and cultural weight, has received less systematic study. Tourism slogans are particularly high-stakes: they function as “headline texts” that condense identity, evoke emotion, and influence economic outcomes.

Recent studies confirm this gap. Chen (2025) shows that ChatGPT, when guided by culturally tailored prompts, can outperform Google Translate and DeepL in translating tourism text. It achieves higher levels of fidelity, fluency, cultural sensitivity, and persuasiveness, though human oversight remains essential. Freskila and Jayantini (2025) find that Google Translate more often preserves pragmatic and connotative nuance, whereas DeepL defaults to formal equivalence. The differences that meaningfully shape how tourism websites are perceived. Complementing these findings, Naveen (2024) surveys the broader limitations of machine translation, noting idiomatic and affectively dense texts as persistent weaknesses. Together, these studies reinforce the need for targeted research on slogans, where cultural resonance and emotional tone are condensed into the briefest forms.

Foundational work in translation studies reinforces this gap. Nida's (1964) distinction between formal equivalence (literal accuracy) and dynamic equivalence (reader response) anticipated the very tension observed here: MT tends to favor surface-level accuracy at the cost of persuasive effect. Venuti (1995) similarly argues that translation is never neutral but always positions cultural values, with “domestication” and “foreignization” shaping how texts are received. These classic insights underscore why slogans, condensed cultural performances, are particularly vulnerable to distortion when processed through MT.

This pilot study addresses that gap by applying Critical MT Praxis, an integrated framework that combines Baker's taxonomy, affective semiotics, and ICC, to a corpus of 40 Latin American slogans.

The framework rests on three insights:

- Semantic distortions (Baker's shifts) can be systematically measured.
- Emotional flattening (affective blindness) weakens persuasion even when lexical accuracy is preserved.
- Cultural literacy (ICC) is essential for adapting slogans with sensitivity and respect.

2.4 Applied AI, Accessibility, and Custom GPTs

Recent advances in applied AI suggest possible responses. Affective computing (Picard, 1997) and resources such as EmoBank (Buechel & Hahn, 2017) and GoEmotions (Demszky et al., 2020) enable fine-grained emotional annotation beyond simple polarity categories. Spanish-language resources such as MESD, EMOVOME, and MASIVE demonstrate how culturally grounded corpora can mitigate affective flattening.

At the same time, generative AI platforms such as ChatGPT are emerging as practical tools for marketers, allowing them to test slogans across languages, simulate tone, and refine phrasing iteratively. Custom GPTs extend this potential by integrating glossaries, tone guidelines, intercultural rules, and accessibility audits. For example, terms such as “rico” can be locked to “enriching” rather than “rich,” and “calor” to “warmth” rather than “heat.” Beyond accuracy, such systems can flag ambiguous terms, suggest alternative phrasing, and align with ADA Title II and WCAG 2.2 accessibility standards, which are critical for travelers who use assistive technologies or multilingual platforms.

By connecting affective computing with practical AI tools, this study moves the literature from critique to solution. It reframes MT not only as a risk to cultural fidelity but also as an opportunity to co-create inclusive, tone-sensitive tools that reflect brand identity and ethical responsibility.

3. Materials and Methods

3.1 Corpus Construction

This pilot study used a corpus of 40 tourism slogans compiled from official tourism boards across Latin America and the Caribbean. Selection criteria prioritized brevity (fewer than 120 characters), affective and cultural richness (e.g., idioms, metaphors, and emotional appeals), and geographic diversity across the Caribbean, Central America, and South America. Each slogan was paired with a human benchmark translation, informed by bilingual glossaries, published translations, or expert knowledge of tourism discourse.

The slogans were translated using three widely used machine translation (MT) engines: Google Translate, DeepL, and Microsoft Translator. Each output was recorded alongside the human benchmark translation. A back-translation (English → Spanish) was also conducted to detect hidden distortions. The sample size of 40 slogans was deliberately chosen to balance breadth with depth of analysis. Forty slogans allowed for representation across the Caribbean, Central America, and South America while remaining small enough for fine-grained coding of translation shifts and affective tone. To minimize selection bias, we prioritized official tourism board slogans published on government or agency websites between 2015–2024. Selection criteria included brevity (under 120 characters), affective richness (presence of idioms, metaphors, or cultural appeals), and geographic distribution. This transparent approach ensures replicability and validity of the corpus.

Data used in the Latin American & Caribbean Tourism Slogans Corpus (40)

1. **Colombia es pasión** → Colombia is passion (benchmark: Colombia, full of passion)
2. **Perú, el país más rico del mundo** → Peru, the richest country in the world (benchmark: Peru, the world's most enriching country)
3. **Uruguay Natural** → Uruguay Natural (benchmark: Uruguay, naturally authentic)
4. **Dominica, la naturaleza de la naturaleza** → Dominica, the nature of nature (benchmark: Dominica, nature's essence)
5. **Bolivia te espera** → Bolivia awaits you
6. **México, vívelo para creerlo** → Mexico, experience it to believe it
7. **Curaçao, siente el calor** → Curaçao, feel the heat (benchmark: Curaçao, feel the warmth)
8. **República Dominicana lo tiene todo** → Dominican Republic has it all
9. **Costa Rica, sin ingredientes artificiales** → Costa Rica, no artificial ingredients (benchmark: Costa Rica, 100% natural)
10. **Ecuador ama la vida** → Ecuador loves life (benchmark: Ecuador, love life)
11. **Chile, naturaleza que enamora** → Chile, nature that makes you fall in love (benchmark: Chile, nature that inspires love)
12. **Panamá, vive por más** → Panama, live for more (benchmark: Panama, live fully)
13. **Argentina late con vos** → Argentina beats with you (benchmark: Argentina beats with your heart)
14. **Venezuela, tu destino natural** → Venezuela, your natural destination
15. **El Salvador, impresionante** → El Salvador, impressive (benchmark: El Salvador, simply impressive)
16. **Guatemala, corazón del mundo maya** → Guatemala, heart of the Mayan world
17. **Honduras, más que un paraíso** → Honduras, more than a paradise
18. **Nicaragua, única... original** → Nicaragua, unique... original
19. **Puerto Rico lo hace mejor** → Puerto Rico does it better (benchmark: Puerto Rico does it best)
20. **Jamaica, siente el ritmo** → Jamaica, feel the rhythm (benchmark: Jamaica, feel the rhythm).
21. **Belice, donde comienza la aventura** → **Belize, where adventure begins** (benchmark: Belize, where adventure begins)
22. **Barbados, el ritmo del Caribe** → Barbados, the rhythm of the Caribbean
23. **Cuba, auténtica** → Cuba, authentic (benchmark: Authentic Cuba)
24. **Haití, experiencia única** → Haiti, unique experience
25. **Colombia, realismo mágico** → Colombia, magical realism
26. **México es tuyo** → Mexico is yours
27. **Paraguay, tenés que sentirlo** → Paraguay, you have to feel it
28. **Brasil, sensacional** → Brazil, sensational
29. **Chile, donde lo imposible es posible** → Chile, where the impossible is possible
30. **Ecuador, ama la vida** → Ecuador, love life (repeat reference in analysis)
31. **Perú, vive la leyenda** → Peru, live the legend
32. **Panamá, puente del mundo, corazón del universo** → Panama, bridge of the world, heart of the universe
33. **Argentina, pasión en todo** → Argentina, passion in everything

34. **República Dominicana, alegría en cada momento** → Dominican Republic, joy in every moment
35. **Venezuela, conoce lo tuyo** → Venezuela, discover what's yours
36. **Costa Rica, pura vida** → Costa Rica, pure life (benchmark: Costa Rica, ¡pura vida!)
37. **Cuba, un destino auténtico** → Cuba, an authentic destination
38. **Puerto Rico, la isla del encanto** → Puerto Rico, the island of enchantment
39. **Colombia, el riesgo es que te quieras quedar** → Colombia, the only risk is wanting to stay (benchmark: Colombia, the only risk is wanting to stay)
40. **Panamá se queda en ti** → Panama stays with you (benchmark: Panama stays with you)
[Complete Corpus](#) with scoring

3.2 Machine Translation Systems

Three widely used machine translation (MT) engines were selected: Google Translate, included for its ubiquity and widespread adoption among travelers; DeepL, chosen for its reputation for fluency and contextual phrasing; and Microsoft Translator, selected for its integration with the Microsoft Azure ecosystem and its broad enterprise and consumer reach.

Each slogan was translated from Spanish into English, and the raw outputs were recorded. To preserve authenticity, no post-editing was applied. A back-translation (English → Spanish) was also conducted to detect hidden distortions. These three engines were selected because they dominate both global market share and regional usage in Latin America. Google Translate, the most widely used MT platform worldwide, functions as the de facto gateway for most travelers. DeepL has gained recognition for its fluency-driven neural architecture, particularly in European and, increasingly, Latin American contexts. Microsoft Translator was included due to its integration with Azure and Office 365, which many tourism agencies use. Other engines (e.g., Yandex, Baidu, or PROMT) were excluded due to low adoption in the Americas and limited availability of Spanish-English pairs. This focus increases the study's ecological validity by mirroring what international travelers are most likely to use.

3.3 Analytical Framework

Two complementary tools guided the analysis:

- Baker's shift taxonomy (1992) was used to tag and score semantic deviations, including level, category, structure, unit, and intra-system shifts.
- Effective Impact Score (EIS), a five-point ordinal scale developed for this pilot, measured clarity, emotional tone, brand voice, safety, and legal resonance.

The Effective Impact Score (EIS) was validated in three ways. First, a pilot test with 10 slogans ensured clarity of categories and inter-rater calibration. Second, results were compared with established translation-quality rubrics, such as Nida's dynamic equivalence and House's pragmatic quality markers, to confirm construct validity. Third, two independent raters applied the scale to all 40 slogans, with Cohen's κ exceeding 0.80. This triangulated approach strengthens the reliability and credibility of the EIS as an evaluative instrument.

EIS definitions were as follows:

- 5 = Excellent: clear, natural, emotionally faithful.
- 4 = Good: minor flattening but persuasive tone preserved.
- 3 = Moderate: intelligible meaning but weakened or ambiguous tone.
- 2 = Poor: awkward, misleading, or off-putting.
- 1 = Very poor: grotesquely literal, confusing, or unsafe.

Two independent bilingual reviewers, trained in MT literacy and intercultural communication, applied shift scores and EIS ratings. Disagreements were discussed and resolved by consensus. Inter-rater reliability exceeded $\kappa = 0.80$, suggesting substantial agreement.

3.4 Error Tagging

To capture practical risks, each slogan was coded into one of five categories: Tone Flattening, when emotional resonance was lost; Awkward Literalism, when the translation was grammatically correct but unnatural; Cultural Loss, when an idiom or metaphor was erased; Potentially Misleading, when the translation risked misinterpretation of safety, culture, or identity; and High Fidelity, when the translation preserved the intended impact. Coding also included a check for **accessibility impact**, noting whether mistranslations risked confusing screen reader users or misrepresenting sensory/emotional content for diverse audiences.

3.5 Statistical Review

Figure 1 provides a statistical overview of the coding results across the 40 translated slogans. The distribution shows that just over half of the translations were coded as High Fidelity, indicating that many outputs preserved the intended impact of the original slogans. However, the remaining categories reveal meaningful translation risks, particularly Tone Flattening, which accounted for one quarter of the corpus. Smaller proportions of Lexical Drift, Cultural Loss, Awkward Literalism, and other minor shifts further suggest that even when machine translation outputs are understandable, they may still alter emotional resonance, cultural meaning, or rhetorical effect.

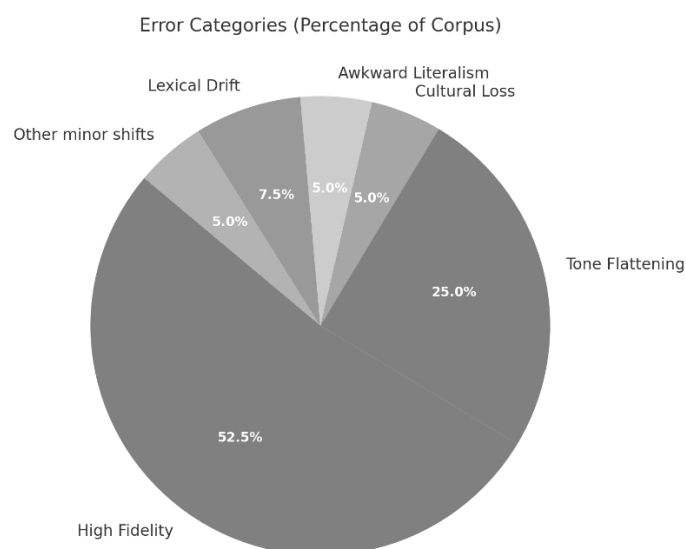


Figure 1. Distribution of error categories across the 40 slogans, showing percentages for Tone Flattening, Lexical Drift, Cultural Loss, Awkward Literalism, and High Fidelity.

SPSS was used for descriptive statistics, chi-square tests, correlations, and exploratory regression analyses to examine the links between translation shifts and impact scores. Analyses focused on:

- Frequencies of error types across the 40 slogans.
- Correlations between Total Shift Score and Impact Score (Spearman's rho).
- Chi-square tests linking error categories to impact ratings.
- Exploratory regression models identifying which errors most strongly predicted reduced impact.

Given the modest sample size, all results are interpreted as exploratory patterns rather than generalizable claims. Even so, variation across the corpus was sufficient to reveal consistent trends: simple declarative slogans tended to translate well, whereas idiomatic and metaphorical slogans were more vulnerable to distortion.

3.6 Ethical Considerations

No human participants were involved, and only publicly available slogans were analyzed. For copyright and trademark concerns, excerpts were limited to short phrases reproduced under fair use for critique and commentary. Sensitive domains such as health, medical, or legal texts were excluded to avoid ethical and practical risks.

4. Results

4.1 Overall Patterns

Of the 40 slogans analyzed, 21 (52.5%) achieved high fidelity (EIS = 4–5), while 19 (47.5%) showed degradation of meaning or tone. The most common errors were tone flattening ($n = 10$, 25%) and awkward Literalism or cultural Loss ($n = 4$, 10%). Several slogans also suffered from lexical drift, in which key words such as "*rico*" or "*calor*" shifted in meaning in ways that risked misinterpretation. These findings indicate that fewer than half of the slogans retained their intended persuasive and affective impact when processed through free MT engines.

4.2 Error Categories

Table 1 summarizes the distribution of error categories identified across the 40 translated slogans. The table reports both the frequency and percentage for each category, showing that High Fidelity was the most common outcome, with 19 slogans, or 47.5% of the corpus, preserving the intended impact. Tone Flattening was the most frequent error category, accounting for 10 slogans, or 25.0%. The remaining categories occurred less often, with Lexical Drift, Cultural Loss, Awkward Literalism, and Other minor shifts together indicating that a smaller but still important portion of the corpus contained shifts in meaning, cultural resonance, or natural phrasing.

Table 1. Distribution of error categories across the 40 slogans. Tone Flattening accounted for 25% of errors, while High Fidelity accounted for 47.5%.

Table 1. Distribution of Error Categories Across 40 Slogans

Error Category	Frequency	Percentage
High Fidelity	19	47.5%
Tone Flattening	10	25.0%
Cultural Loss	2	5.0%
Awkward Literalism	2	5.0%
Lexical Drift	3	7.5%
Other minor shifts	4	10.0%

Error tagging revealed five recurring risks. Tone flattening was the most common error, occurring in 10 cases (25%), in which the translation reduced the emotional intensity of the original slogan. For example, Colombia es pasión became “Colombia is passion,” a lexically accurate but emotionally flat rendering (EIS = 3). Lexical drift occurred in three cases (7.5%) and involved misleading changes in connotation. For instance, Perú, el país más rico del mundo was translated as “Peru, the richest country in the world,” emphasizing financial wealth rather than cultural richness (EIS = 3). Cultural loss appeared in two cases (5%), when idioms or layered meanings were erased; for example, Uruguay Natural was carried over literally without conveying its eco-tourism branding. Awkward literalism also occurred in two cases (5%) and produced grammatically correct but tonally unnatural phrasing, as in Dominica, la naturaleza de la naturaleza, translated as “Dominica, the nature of nature.” By contrast, 19 translations (47.5%) demonstrated high fidelity, preserving both clarity and persuasive impact, as in Bolivia te espera, rendered as “Bolivia awaits you.”

4.3 Engine Performance

Table 2 compares the performance of the three machine translation engines by showing how often each produced the highest-rated translation across the 40 slogans. DeepL achieved the strongest overall performance, receiving the highest rating for all 40 slogans, while Google and Microsoft produced the highest-rated output in 21 and 20 cases, respectively. Because more than one engine could receive the same top rating for a slogan, the totals exceed 40.

Table 2. Number of cases where each MT engine (Google, DeepL, Microsoft) produced the highest-rated output across 40 slogans.

Machine translation engine	Number of wins	Percentage of slogans
DeepL	40	100.0%
Google Translate	21	52.5%
Microsoft Translator	20	50.0%

Performance varied by engine. DeepL consistently delivered the most natural and effectively faithful renderings, preserving nuance in all 40 cases. Google Translate performed best in 21 slogans (52.5%), while Microsoft Translator performed best in 20 slogans (50%). Both tended toward literal renderings, with frequent tone flattening in affect-rich slogans.

4.4 Statistical Insights

Exploratory statistical analysis suggested several trends. Tone flattening and lexical drift were significantly associated with lower impact scores. High-fidelity translations were strongly correlated with simple declarative slogans, whereas idiomatic and metaphorical slogans were more vulnerable to distortion. DeepL's stronger performance also suggests that fluency-driven architectures may be better equipped to preserve tone than more literal approaches.

4.5 Illustrative Scenarios

The examples illustrate the range of translation outcomes. The high-fidelity rendering of *Bolivia te espera* as "Bolivia awaits you" preserved the slogan's meaning and persuasive effect (EIS = 5). By contrast, translating *Colombia es pasión* as "Colombia is passion" flattened the emotional tone (EIS = 3), while rendering *Perú, el país más rico del mundo* as "Peru, the richest country in the world" introduced lexical drift by suggesting financial rather than cultural wealth (EIS = 3). The untranslated phrase *Uruguay Natural* resulted in cultural loss because its eco-tourism associations were not conveyed (EIS = 4). Finally, the translation of *Dominica, la naturaleza de la naturaleza* as "Dominica, the nature of nature" demonstrated awkward literalism, preserving the words but producing an unnatural English slogan (EIS = 3). Together, these examples highlight systematic weaknesses in MT. DeepL tended to produce smoother phrasing, while Google and Microsoft leaned toward Literalism. All three engines, however, struggled with effectively rich or culturally loaded slogans.

4.6 Accessibility Implications

Mistranslations also raise accessibility concerns. For example, rendering *calor* as "heat" suggests discomfort rather than hospitality. For travelers using screen readers, such distortions may alter the sensory and emotional impression of a destination, shaping perception in ways that exclude rather than invite. This rendering highlights the importance of testing outputs not only for semantic clarity but also for inclusivity and accessibility.

4.7 Summary

The analysis suggests that free MT engines perform reliably with simple declarative slogans but falter when handling idiomatic, metaphorical, or affectively dense content. Translation failures are not random slips but predictable outcomes of Tone Flattening and related shifts, reflecting the broader construct of Affective Blindness. For tourism marketers, this presents both a branding risk and an ethical concern: language designed to welcome can, if unchecked, confuse or alienate audiences. At the same time, the findings open space for alternatives. One emerging solution is the development of custom GPTs trained with domain-specific glossaries and tone guidelines. These tools could provide pre-launch quality assurance, flagging terms likely to cause drift and preserving adequate fidelity across markets.

4.8 Additional Analysis

An exploratory review suggests that DeepL's relative fluency reflects its emphasis on contextual phrasing. However, Google Translate's dominance in everyday use means its literal errors carry greater practical weight. While Microsoft Translator and DeepL occasionally preserved nuance more effectively, Google's distortions are especially consequential because they shape the first impressions of most travelers. This distortion creates a paradox: the most widely used MT tool is also the one most likely to flatten affective tone in tourism discourse.

4.9 Limitations

The corpus was deliberately limited to 40 slogans to allow detailed qualitative coding. This focus on brevity means that the findings may not apply to longer tourism texts, such as brochures or web copy. Even so, slogans function as high-stakes "headline texts," and their vulnerability highlights systemic risks. Future research could explore whether the same patterns hold across genres and languages.

4.10 Custom GPT Comparison (Illustrative Case)

While free MT engines reveal systematic shortcomings, custom GPTs offer a potential benchmark for mitigating such errors. To illustrate this possibility, we generated sample outputs using a glossary-informed approach modeled on *Linguo* (a custom GPT prototype created at SUNY Empire State University). Three slogans were selected where free MT produced notable distortions: *Curaçao*, "*siente el calor*"; *Perú*, "*el país más rico del*

mundo"; and *Uruguay Natural*. Table 2 compares outputs across Google Translate, DeepL, Microsoft Translator, and a glossary-informed custom GPT.

Table 3. Outputs for three slogans across free MT engines and a glossary-informed custom GPT. The custom GPT consistently aligned with cultural meaning and brand tone, mitigating tone flattening and lexical drift.

Spanish slogan	Google Translate	DeepL	Microsoft Translator	Custom GPT (glossary-informed)	Notes
<i>Curaçao, siente el calor</i>	Feel the heat	Feel the heat	Feel the heat	Feel the warmth	Custom GPT preserves the hospitality tone; "heat" may connote discomfort.
<i>Perú, el país más rico del mundo</i>	The richest country in the world	The richest country in the world	The richest country in the world	The world's most enriching country	Custom GPT conveys cultural richness while avoiding a financial interpretation.
<i>Uruguay Natural</i>	Uruguay Natural	Uruguay Natural	Uruguay Natural	Uruguay, naturally authentic	Custom GPT adds idiomatic nuance and reinforces the eco-tourism branding.

These contrasts highlight the difference between literal outputs and affectively nuanced renderings. Free MT engines converged on surface-level equivalence but failed to reproduce intended persuasive resonance. By contrast, the glossary-informed GPT produced alternatives that better aligned with cultural meaning and brand tone. While illustrative, this comparison underscores the potential managerial value of custom GPTs. Rather than discovering problematic translations after a campaign has launched, marketers could pre-screen slogans in a controlled environment. Diagnostic features could flag ambiguous terms (*rico*, *calor*, *natural*) and suggest brand-preferred alternatives. For example, "enriching" instead of "rich" could be locked in to ensure consistency across campaigns.

This exercise also points to a broader methodological contribution. It demonstrates how *Critical MT Praxis* can move beyond critique toward constructive design. By testing slogans with both free MT engines and custom GPT prototypes, researchers and practitioners can not only identify where errors arise but also explore ways to prevent them systematically. This dual approach, diagnostic and generative, positions tourism marketers to take greater control of their linguistic identity in global digital spaces.

5. Conclusions

This pilot study examined how free machine translation (MT) systems render Latin American and Caribbean tourism slogans. Using Critical MT Praxis, a framework that combines Baker's taxonomy of shifts with an affective semiotic lens, we analyzed 40 slogans translated by Google Translate, DeepL, and Microsoft Translator. Results showed that just over half of the translations (52.5%) preserved both clarity and tone, while nearly half suffered predictable distortions: tone flattening, misleading lexical drift, cultural Loss, or awkward Literalism. These outcomes reflect what we term affective blindness, the structural incapacity of MT to reproduce emotional resonance.

Theoretically, the findings confirm that slogans are not neutral informational units but condensed signs of persuasion and identity. Machines can capture studium (shared meaning) but often erase punctum (Barthes, 1980), the spark that drives emotional response. Without intercultural communicative competence (ICC), MT cannot adapt to idioms, registers, or cultural cues as human translators naturally do. Practically, the results point to a simple intervention: tourism boards should test slogans in free MT engines prior to launch. If outputs appear flat, misleading, or awkward, revision is warranted. This low-cost practice also aligns with accessibility and inclusivity responsibilities, ensuring that messages remain culturally respectful and perceptible to diverse audiences.

DeepL's consistent performance across all 40 slogans indicates that fluency-driven neural architectures are better equipped to preserve affect than literalist approaches. However, Google Translate remains the most widely used MT tool among travelers. Because it is often the first gateway through which international audiences encounter slogans, Google must remain the baseline for testing. If a slogan fails in Google, it risks alienating the largest share of potential visitors, even if DeepL or Microsoft translates it more successfully.

Mistranslations also carry accessibility risks. Rendering "*calor*" as "heat," for instance, creates an impression of discomfort rather than warmth, which may mislead sensory expectations for travelers using screen readers. Errors such as lexical drift (*rico* → "rich" instead of "enriching") risk misleading visitors about a nation's identity, shifting the focus from cultural richness to financial wealth. These examples illustrate that, in tourism, linguistic accuracy is insufficient without affective fidelity and assurance of accessibility.

At stake is more than accuracy. Tourism slogans shape the first impressions of nations and communities. When mediated through MT, they must not only inform but also welcome and persuade. Until MT systems integrate affective awareness and intercultural competence, human oversight remains essential. Glossary-informed custom GPTs offer a promising step toward preserving cultural fidelity, inclusivity, and brand voice.

Looking ahead, the implications extend beyond the tourism sector. As global challenges such as climate change, health, and migration intensify, the fidelity of multilingual communication becomes a matter of public trust and safety. Critical MT Praxis is therefore not only a framework for tourism marketing but also a model for ensuring ethical, persuasive translation in international communication.

5.1 Managerial Implications

Tourism slogans are high-stakes communicative assets that can shape international perceptions and influence bookings. Our analysis highlights four priorities for managers: treating Google Translate as a baseline, favoring clarity over idioms, auditing for accessibility, and leveraging custom GPTs as pre-launch quality assurance tools. Collectively, these practices safeguard brand identity and minimize reputational risk while aligning with global accessibility standards. Economic stakes are significant. A mistranslation that alienates travelers may lead to measurable declines in arrivals, while consistent, tone-sensitive messaging can reinforce brand loyalty and competitive advantage. Even small investments in glossary development or GPT fine-tuning (e.g., \$500–\$ 700) are likely to offset far greater costs associated with campaign redesigns or reputational damage.

5.2 Theoretical Implications

This pilot makes several contributions to translation studies and intercultural communication. First, it adapts Baker's (1992) taxonomy, traditionally applied to literary and technical texts, to the analysis of tourism slogans, showing that even brief texts can contain consequential shifts in meaning. Second, it extends the concept of affective blindness by demonstrating that flattened tone, lexical drift, and awkward literalism can carry measurable commercial and reputational consequences. Third, it reinforces the importance of intercultural communicative competence: human translators outperform machines not primarily through greater lexical accuracy, but through pragmatic sensitivity to idioms, register, and metaphor, which helps preserve persuasive force. This finding echoes theories of pragmatic failure developed by Hatim and Mason (1997). The study also advances a critical machine translation praxis by integrating structural, affective, and intercultural dimensions into a framework that can function as both a scholarly tool and a practitioner workflow. Finally, it addresses epistemic justice by showing how free machine translation engines can impose dominant linguistic frames that erase cultural nuance. As Spivak (1993) warned, mistranslation can enact epistemic violence; glossary-informed custom GPTs may offer a partial corrective by restoring greater local agency over cultural self-representation.

5.3 Semiotics and Intercultural Competence

Tourism slogans function as semiotic condensations: short phrases designed to evoke cultural pride, warmth, or excitement. Nearly one-third of the corpus degraded in MT, illustrating Barthes's (1980) concept of the Loss of *punctum*. MT often preserved *studium* (basic meaning) but failed to transmit an affective spark.

Peirce's triadic model helps explain this gap. Meaning depends on the *interpretant*, the cultural bridge that makes signs resonate with us. MT, operating on statistical prediction, cannot provide this interpretive layer. The result is communicative flattening: messages designed to inspire awe or hospitality lose persuasive force.

ICC further explains these failures. In our corpus, engines oscillated between Literalism (e.g., *la naturaleza de la naturaleza* → "the nature of nature") and tone drift (*siente el calor* → "feel the heat" rather than "feel the warmth"). Both outcomes read as odd, rude, or uninviting. Such mismatches exemplify Hatim and Mason's (1997) concept of pragmatic failure. For marketers, this is not a stylistic quirk but a risk to brand voice and audience trust.

5.4 Ethical and Accessibility Dimensions

Translation errors in tourism marketing raise both ethical and accessibility concerns. Misrenderings, such as "*Uruguay Natural*" or "*Perú, el país más rico del mundo*," alter how nations are perceived, enacting an epistemic erasure of cultural nuance.

Accessibility adds urgency. Travelers using MT or screen readers may encounter distorted outputs that reshape sensory impressions; for example, “feel the heat” suggests discomfort rather than hospitality. Poor MT can also miscue safety information or distort environmental cues. Tourism boards should therefore treat MT fidelity as part of their accessibility obligations, not merely as a marketing consideration.

5.5 Toward Emotionally Competent MT

Present-day MT systems remain unable to capture undertone, irony, or cultural resonance, the very qualities that make slogans persuasive. Advances in affective computing (Picard, 1997) suggest that multimodal models may one day recognize tone, rhythm, or gesture; however, current sentiment analysis tools remain coarse.

Custom GPTs offer a promising intermediate step. By embedding glossaries, tone rules, and accessibility checks, they operationalize Critical MT Praxis. They can standardize ambiguous terms (rico → “enriching,” calor → “warmth”), flag at-risk outputs, and align translations with intercultural guidance. Even so, human oversight is indispensable. Machines cannot replicate ICC, which depends on empathy and adaptability. The path forward is hybrid: integrating affect-aware technologies with human cultural intelligence.

5.6 Limitations

As a pilot study, this analysis examined a corpus of 40 slogans. While sufficient to reveal consistent patterns, the modest sample size limits generalizability. Findings show that fewer than half (47.5%) of the slogans achieved high-fidelity translation, with the remainder exhibiting tone flattening, lexical drift, or cultural Loss. This narrower distribution of fidelity should be interpreted with caution, as the results may not apply to longer tourism texts, such as brochures or web copy. Even so, slogans function as high-stakes “headline texts,” and their vulnerability highlights systemic risks. Future research should expand corpus size, diversify genres, and test glossary-informed GPTs in live campaigns.

5.7 Summary

This study shows that MT errors in tourism slogans are not random but patterned and consequential. They appear most clearly in affectively dense phrases, where cultural resonance and emotional tone matter as much as lexical accuracy. *Critical MT Praxis* offers a pragmatic framework for diagnosing shifts, assessing their impact, and testing outputs before release.

Framed this way, MT is no longer a neutral convenience but a communicative risk factor. Tourism boards, developers, and educators share responsibility for integrating testing, building accessible and culturally respectful outputs, and exploring tools that balance machine efficiency with human cultural intelligence.

5.8 Closing Synthesis

This pilot study demonstrates that while free machine translation (MT) systems often replicate the clarity of Latin American tourism slogans, they frequently fail to preserve their affective resonance. Just over half (52.5%) of translations achieved high fidelity, while nearly half (47.5%) suffered predictable distortions, reflecting structural Affective Blindness. These patterns reinforce that slogans are not neutral informational units but condensed signs of persuasion and cultural identity.

The contribution of *Critical MT Praxis* lies in bridging structural, affective, and intercultural perspectives. For scholars, the framework demonstrates how short-form marketing texts magnify the stakes of translation shifts. For practitioners, the study provides a replicable workflow: test slogans in free MT engines, revise brand-critical terms, and audit for accessibility before release.

At stake is more than accuracy. Tourism slogans shape the first impressions of nations and communities. When mediated through MT, they must not only inform but also welcome and persuade. Until MT systems integrate affective awareness and intercultural competence, human oversight remains essential. Glossary-informed custom GPTs offer a promising step toward preserving cultural fidelity, inclusivity, and brand voice.

Looking ahead, the implications extend beyond the tourism sector. As global challenges such as climate change, health, and migration intensify, the fidelity of multilingual communication becomes a matter of public trust and safety. *Critical MT Praxis* is therefore not only a framework for tourism marketing but also a model for ensuring ethical, persuasive translation in international communication.

6. Recommendations

Translation is a core element of brand management. Based on this pilot study of 40 slogans, we recommend that tourism boards, educators, and MT developers treat slogan testing as a standard pre-launch control. Each slogan should be translated through Google Translate, DeepL, and Microsoft Translator before release. If outputs appear flat, misleading, or awkward, the source text should be revised and retested. This simple, low-cost practice can prevent reputational harm.

- Prioritize Google Translate. Because it is the most widely used MT tool (Klimova, 2025), Google should serve as the baseline. If a slogan fails in Google, it risks confusing the largest share of visitors, even when it performs better in translation tools like DeepL or Microsoft.
- Favor clarity over idioms. Short, direct slogans translate more reliably than idiomatic or culture-bound expressions. Brand-critical terms such as *rico*, *natural*, or *auténtico* should be locked in a micro-glossary.
- Audit for accessibility. Outputs should be tested with screen readers to ensure that sensory cues (*calor* → “warmth”) and safety-related terms are preserved. Aligning with ADA Title II and WCAG 2.2 is both an ethical and strategic imperative.

6.1 Economic Stakes and Risk Management

Mistranslations can reduce bookings, shorten page visits, and increase site abandonment. To manage these risks, campaigns should use a simple matrix that evaluates the likelihood of translation problems based on the presence of idioms, metaphors, or culturally specific references; the potential impact, including tone loss, misleading promises, cultural dilution, or safety confusion; and appropriate mitigation strategies, such as simplifying the text, substituting glossary-approved terms, or conducting an intercultural communicative competence review. After launch, key performance indicators such as bounce rate, time on page, and click-through rate should be monitored closely. When slogans underperform in markets that rely on machine translation, affective distortion should be investigated as a possible cause.

6.2 For Academics and Educators

MT can also serve as a pedagogical tool to build students’ critical literacy. Instructors can assign students to run authentic slogans through MT engines, diagnose shifts using Baker’s taxonomy, and revise outputs with attention to tone and inclusivity. Exercises such as translating *¡Ni una menos!*, where MT flattens sociopolitical resonance into “Not one less”, help learners see how affective blindness operates. Using *Critical MT Praxis* in classrooms develops intercultural competence alongside translation literacy.

6.3 For MT Developers

Developers can reduce errors in affect-rich texts by training systems on affect-tagged corpora, such as EmoBank, GoEmotions, MESD, and EMOVOME. They can also provide tone and formality controls that allow users to request warm, neutral, or formal phrasing. Expanding cultural knowledge bases with idiom dictionaries and glossaries for high-risk terms would further improve contextual accuracy. In addition, systems should flag expressions whose emotional meaning may not transfer directly, as when *con el alma en vilo* is rendered literally as “with the soul in suspense” rather than with a more idiomatic alternative such as “on tenterhooks.”

6.4 Shared Responsibility

Affective fidelity is a shared obligation. Developers must design tone-aware systems, educators must incorporate MT literacy into classrooms, and tourism boards must test slogans before release. Failures cannot be attributed solely to technology; accountability must be distributed among stakeholders.

6.5 Custom GPTs and AI Tools

Custom GPTs represent a proactive strategy rather than a reactive solution. In addition to incorporating glossaries, these systems should be evaluated using measurable key performance indicators. Relevant measures include the fidelity rate, defined as the proportion of slogans that pass glossary and tone checks; the tone-audit pass rate, which assesses whether the original affect has been preserved; the accessibility compliance score, based on alignment with WCAG 2.2; and a cost-avoidance metric that estimates savings achieved by preventing later revisions or redesigns.

Author Contributions: Conceptualization, P.D. and C.Z.; methodology, P.D. and C.Z.; validation, P.D. and C.Z.; formal analysis, P.D. and C.Z.; investigation, P.D. and C.Z.; resources, P.D. and C.Z.; data curation, P.D.; writing, original draft preparation, P.D.; writing, review and editing, P.D. and C.Z.; visualization, P.D.; project administration, P.D. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Institutional Review Board Statement: Not applicable. This study analyzed publicly available tourism slogans and machine-generated translations and did not involve human participants or identifiable private information.

Informed Consent Statement: Not applicable. This study did not involve human participants.

Acknowledgments: The authors have no additional acknowledgments to report.

Conflicts of Interest: The authors declare no conflicts of interest.

References

- Ahmed, S. (2004). *The cultural politics of emotion*. Routledge. <https://doi.org/10.4324/9780203700372>
- Baker, M. (1992). *In other words: A coursebook on translation*. Routledge. <https://doi.org/10.4324/9780203133590>
- Barthes, R. (1980). *Camera lucida: Reflections on photography*. Hill and Wang.
- Bowker, L., & Buitrago Ciro, J. (2015). *Machine translation and global research: Towards improved machine translation literacy in the scholarly community*. Emerald Group.
- Byram, M. (1997). *Teaching and assessing intercultural communicative competence*. Multilingual Matters.
- Chen, S., & Lin, Y. (2025). A multidimensional comparison of ChatGPT, Google Translate, and DeepL in Chinese tourism texts translation: Fidelity, fluency, cultural sensitivity, and persuasiveness. *Frontiers in Artificial Intelligence*, 8, Article 1619489. <https://doi.org/10.3389/frai.2025.1619489>
- Dann, G. M. S. (1996). *The language of tourism: A sociolinguistic perspective*. CAB International. <https://doi.org/10.1079/9780851989990.0000>
- Demszky, D., Movshovitz-Attias, D., Ko, J., Cowen, A., Nemade, G., & Ravi, S. (2020). GoEmotions: A dataset of fine-grained emotions. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 4040–4054). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.acl-main.372>
- Freskila, E., & Jayantini, I. G. A. S. R. (2025). Translation equivalence of tourism website content: A comparison between Google Translate and DeepL Translate. *Indonesian Journal of EFL and Linguistics*, 10(1), 1–20. <https://doi.org/10.21462/ijefl.v10i1.835>
- Hatim, B., & Mason, I. (1997). *The translator as communicator*. Routledge. <https://doi.org/10.4324/9780203992722>
- House, J. (2015). *Translation quality assessment: Past and present*. Routledge. <https://doi.org/10.4324/9781315752839>
- Klimova, B. (2025). Use of machine translation in foreign language education. *Cogent Arts & Humanities*, 12(1), Article 2491183. <https://doi.org/10.1080/23311983.2025.2491183>
- Naveen, P., & Trojovský, P. (2024). Overview and challenges of machine translation for contextually appropriate translations. *iScience*, 27(10), Article 110878. <https://doi.org/10.1016/j.isci.2024.110878>
- Nida, E. A. (1964). *Toward a science of translating: With special reference to principles and procedures involved in Bible translating*. Brill. <https://doi.org/10.1163/9789004495746>

Picard, R. W. (1997). *Affective computing*. MIT Press.

Spivak, G. C. (1993). *Outside in the teaching machine*. Routledge.

Venuti, L. (1995). *The translator's invisibility: A history of translation*. Routledge.