

Pragmatic Failures in Machine Translation of Cross-border Comments: An Intercultural Perspective

Mingzhong Jiang*

Tianjin University of Finance and Economics, Tianjin 300202, China

**Author to whom correspondence should be addressed.*

Copyright: © 2026 Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY 4.0), permitting distribution and reproduction in any medium, provided the original work is cited.

Abstract: From the perspective of cross-cultural communication, this study analyzes pragmatic failures in machine translation (MT) of cross-border user comments on TikTok and Xiaohongshu. Based on Thomas's (1983) pragmatic failure theory, we compare original comments with MT outputs (Google Translate, DeepL) and human translations. Pragmalinguistic failures show up as literal or odd renditions of slang, irony, and internet neologisms. Sociopragmatic failures come from MT ignoring cultural taboos, face strategies, politeness norms, and humor—e.g., turning a Chinese self-deprecating expression into an embarrassingly direct English statement, which can cause negative emotions or stereotyping. Such failures hurt genuine cross-cultural understanding and may even trigger unintended offense. The study suggests boosting MT's cross-cultural adaptability via culture-aware corpora and human-in-the-loop post-editing. This research offers empirical evidence for translation evaluation and cross-cultural communication education in the social media era.

Keywords: Machine translation; pragmatic failure; cross-cultural communication; social media comments

Online publication: February 26, 2026

1. Introduction

With globalization and digitalization, social media platforms like TikTok and Xiaohongshu have become vital arenas for cross-border user interaction. Cross-border comments on these platforms predominantly rely on machine translation for interlanguage communication, but this often shows biases when processing culturally specific pragmatic information, leading to pragmatic failures. Based on Thomas's (1983) pragmatic failure theory, such failures can be categorized into pragmalinguistic and sociopragmatic ones. Current research primarily focuses on syntactic or semantic failures^[1], with insufficient attention to the pragmatic dimensions of machine-translated content in social media comments. This study uses cross-border comments from TikTok and Xiaohongshu as corpus material, comparing machine translation with human translation to systematically analyze the types, causes, and impacts of pragmatic failures. It also explores improvement

strategies for enhancing machine translation’s cross-cultural adaptability, providing empirical evidence for translation quality assessment and cross-cultural communication education.

2. Core Concepts Definition

2.1. Classification and Connotation of Pragmatic Failure

Pragmatic failures occur in cross-cultural communication when the intended implications of discourse are misunderstood, leading to deviations from expected outcomes. British linguist Thomas (1983) divided them into two types. One is pragmalinguistic failures, where linguistic forms mismatch their functions. For example, we can see misunderstandings caused by negative transfer of native expressions. The other is sociopragmatic failures, where the original intent is perceived as offensive because of cognitive differences in sociocultural norms or power dynamics. In machine translation, pragmatic failures stand out more because systems rely on parallel corpora and statistical models but don’t really grasp context, tone, or implicit cultural meanings. For irony, self-deprecation, puns, internet neologisms, or culturally specific polite expressions, MT often gives literal or twisted outputs. Both error types can appear together. This can block accurate information flow and may also trigger negative emotions or stereotypes among cross-cultural users. This acts as an implicit barrier to international communication^[2].

2.2. The History of Machine Translation

Machine translation has moved from rule-based and statistical methods to neural models. Neural Machine Translation (NMT) has greatly improved fluency and grammar. For comment translation, NMT has become a key tool for user interaction thanks to its real-time processing, low cost, and multilingual support. But social media comments use colloquial language, fragmented content, heavy context dependence, and lots of internet slang, emojis, irony, and puns—these pose real challenges for MT. While NMT does well on syntax and word choice, its reliance on statistical patterns limits its grasp of deeper intentions, cultural context, and emotional nuance, which can amplify biases in training data^[3]. It often takes euphemisms or humor as blunt statements. Studies show most failures happen at the pragmatic level, e.g., missing social functions like modesty, face-saving, or playful remarks. So evaluating MT from a cross-cultural communication angle needs not just linguistic accuracy but also its ability to deliver appropriate pragmatic effects across different cultures.

2.3. Intercultural Communication

Intercultural communication is when people from different cultural backgrounds exchange information and meaning, which calls for deep adjustments in values, social norms, and politeness principles. In the social media era, such communication is immediate, anonymous, and based on weak ties. Relevant studies show that social distance and group acceptance are key to how well it works^[4]. Good intercultural communication requires understanding not just literal meanings but also underlying attitudes, emotions, and social intentions. Though today’s MT is highly efficient, it can’t mimic human contextual reasoning or cultural empathy. From an intercultural communication standpoint, the key question is: Can machines help people have meaningful, non-misleading cross-language interactions while respecting cultural differences? That’s where this study starts.

2.4. Cross-border Comments

The recent arrival of “foreign Douyin refugees” on cross-border platforms like Xiaohongshu shows how urgent and complex social media comment translation has become^[5]. Pragmatic failures in MT aren’t

accidental—they come from systematic neglect of context and cultural norms. Pragmalinguistic failures show up as mechanical literal translations of slang, irony, and internet neologisms, which undermine self-deprecating intentions. Sociopragmatic failures make it harder to spot cross-cultural taboos, face-saving strategies, and etiquette norms, often producing offensive translations. Such failures block emotional communication and may deepen cultural misunderstandings or even conflicts. So simple algorithm tweaks aren’t enough. We need corpora with culturally sensitive information, human-machine collaborative post-editing, and a “translation feedback” feature that lets users flag bad translations.

3. Research Design

3.1. Data Source Selection Criteria

This study used bilingual (Chinese-English) user reviews from TikTok and Xiaohongshu between 2023 and 2025 as the corpus source. To keep data representative and systematic, we set up multi-tiered screening criteria.

First, platform and topic selection. We picked TikTok and Xiaohongshu because they are among the most used social media platforms for cross-border user interaction, with their built-in translation features widely used. On TikTok, we chose highly interactive tags like #funny, #tutorial, and #reaction—tags that easily produce non-literal readings; on Xiaohongshu, we focused on active topics among Chinese users, including #daily life, #complaints, and #learning.

Second, comment screening criteria. Each comment must meet one of these conditions: (a) Clearly show pragmatic features that might cause misunderstanding, like self-mockery, irony, exaggeration, or internet slang; (b) The reply includes cross-cultural user interactions, hinting at possible cross-border communication cues; (c) The comment itself contains culture-specific elements.

Third, data cleaning and sampling. We first collected about 800 comments. Then we removed pure emoticons, meaningless replies, ads, and content with personal privacy info, leaving 450 valid entries. To keep sample balance, we used stratified random sampling to pick 200 from the 450, making sure 100 comments in Chinese and 100 in English, with roughly equal numbers of four subcategories—online neologisms, culture-specific terms, politeness strategies, and humorous expressions—in each language.

Fourth, the corpus annotation process. We compiled each comment’s original text, Google Translate outputs, DeepL outputs, and human reference translations into a parallel corpus. Two researchers with cross-cultural communication expertise independently annotated three dimensions: (1) presence of pragmatic failures; (2) failure types; and (3) specific manifestations of failures. For conflicting cases, a third expert was consulted to reach consensus. **Table 1** shows the structural distribution of the final corpus.

Table 1. Distribution of Corpus Sample Structures

Terrace	Language	New internet slang	Cultural specificity	Politeness strategies	Humorous expression	Subtotal
TikTok	English	12	13	12	13	50
	the Chinese language	13	12	13	12	50
Little Red Book	English	13	12	13	12	50
	the Chinese language	12	13	12	13	50
Total		50	50	50	50	200

3.2. Framework Analysis

This paper will use Thomas's (1983) pragmatic failure theory and divides failures in machine translation into pragmalinguistic and sociopragmatic. In analysis, we compare machine-translated and human-translated versions of each comment to identify failure types, and explain possible user cognitive biases in cross-cultural interaction contexts. The research also looks at similarities and differences in failure patterns across various MT systems, showing both common weaknesses and individual variations among mainstream systems at the cross-cultural pragmatic level.

3.3. Comparison Method

This study uses a three-way comparison: for each selected social media comment, the original text, machine translation output, and human reference translation are shown side by side. By looking at semantic gaps and pragmatic effects among these three versions, we identify and categorize pragmatic failures. The comparison has three steps. First, two researchers independently compare the machine and human versions, spotting sentence parts with meaning shifts, odd tones, or loss of cultural information. Second, given the interactive context when the comment was posted, they decide if such shifts count as pragmalinguistic or sociopragmatic failures. Third, we compare Google Translate and DeepL on the same source text to measure how often each failure type occurs and in what patterns. To reduce subjective bias, cases where judgments differ are discussed with a third scholar in cross-cultural communication to reach agreement. Through this method, the study shows the pragmatic weaknesses and common challenges that mainstream MT systems face when handling cross-border social media comments.

4. Types and Quantitative Analysis of Pragmatic Failures

4.1. Overall Distribution Frequency

Among the 200 original comments (100 in Chinese and 100 in English), each comment was translated separately using Google Translate and DeepL, resulting in 200 translated outputs from each machine translation system. Statistical analysis identified a total of 312 pragmatic failures (multiple failures may occur within a single translation). **Table 2** presents the frequency distribution by failure type and machine translation system.

Table 2. Frequency Distribution of Two Types of Pragmatic Failures in Machine Translation Output

Failure Type	Google Translate(N = 200)	DeepL (N = 200)	χ^2 price	p price
Pragmatic linguistic failure	78(39.0%)	65(32.5%)	1.84	0.175
Social pragmatic failure	96(48.0%)	73(36.5%)	5.42	0.020*
No failure	26(13.0%)	62(31.0%)	18.88	< 0.001***

The table shows that the differences between Google Translate and DeepL in pragmatic language failures are not significant ($\chi^2 = 1.84$, $p = 0.175$); the differences in social pragmatic failures are significant ($\chi^2 = 5.42$, $p = 0.020$); and the differences in failure-free rates are highly significant ($\chi^2 = 18.88$, $p < 0.001$). Overall, DeepL achieves a higher failure-free rate than Google Translate, indicating its more stable pragmatic processing performance in this sample.

4.2. Pragmalinguistic Failure

4.2.1. Literal Translation and Semantic Shift of Internet Neologisms

The literal or inaccurate translation of internet buzzwords stems from machine translation’s failure to capture the dynamic socio-cultural context and metaphorical meanings, resulting in the complete loss of these neologisms’ inherent pragmatic functions during cross-language conversion. A user commented below a video on Xiaohongshu.

Table 3. Translation Comparison of Misunderstandings Between “Neijuan” and “Po Fang”

addresser	text	Google Translation	DeepL Translation	Correct translation
user	This blogger is just too competitive—I’m completely overwhelmed.	This blogger is so enthusiastic! I’m totally smitten.	The blogger is so rolled, I am broken.	The blogger is way too competitive. I’m totally overwhelmed .

In **Figure 1**, Google Translate translated “卷” as “enthusiastic” and “破防” as “smitten” —both positive or neutral translations that starkly contrast with the original review’s negative pragmatic intentions such as anxiety, helplessness, and self-mockery. This literally “correct” translation precisely constitutes a typical pragmatic linguistic failure.

This case has two translation failures, both pragmalinguistic. Online, “juan” means excessive irrational internal competition or self-sacrifice; Google’s “enthusiastic” turns it into a positive trait, while DeepL’s “so rolled” is literal and loses the social metaphor. “Pofang” started as gaming slang for emotional breakdown under too much stress, often used with self-deprecating humor. Google translated it as “smitten,” which means genuine affection, while DeepL’s “broken” suggests total collapse, not a humorous or exaggerated breakdown. These mistranslations twist the lighthearted, self-deprecating tone of the original comments and can even make readers get the opposite meaning.

4.2.2. Misrecognition of ironic intent

Irony often uses positive words to criticize or mock, but machine translation can’t catch non-literal meanings and just translates word for word, so the original intent gets completely lost.

Table 4. Ironic comments under low-quality tutorial videos on TikTok

addresser	text	Google Translation	DeepL Translation	Manual Translation (Reference)
user	Wow, that method is truly brilliant—no wonder no one thought of it.	Wow, that’s such a clever idea! No wonder nobody thought of it.	Wow, you are so smart with this method, no wonder nobody came up with it.	Wow, what a brilliant method you’ve got there. No wonder no one else thought of it.

Table 4 shows both Google Translate and DeepL turn ironic comments into positive statements through literal translation, losing the sarcastic tone. To back this up, Figure 2 shows an actual Google Translate output. This literalization isn’t random. It reflects how current MT systems lack pragmatic recognition.

In this case, both Google Translate and DeepL failed to detect the ironic intent in expressions such as “too smart” and “no wonder no one thought of it,” instead translating them literally as sincere praise and thereby completely losing the original text’s critical and satirical tone toward the low-quality tutorial. This constitutes a classic pragmatic linguistic failure: the machine only processed the literal meanings of positive terms but failed to decode the negative implications behind the double entendre. Target language readers

may mistakenly interpret the user’s words as genuine praise, leading to misinterpretations in cross-cultural communication.

4.2.3. Semantic Loss in Abbreviations and Emojis

Emoticons play a crucial pragmatic regulatory role in cross-cultural communication on social media, yet their meanings may vary significantly across different cultural groups^[6]. The following example illustrates typical failures in machine translation when processing composite pragmatic signals consisting of “emoticons + text.”

Table 5. Comparison of translations between the abbreviation “lol” and the emoji “speechless”

addresser	text	Google Translation	DeepL Translation	Manual Translation
user	Your singing is so good lol [speechless]	You’re singing so well— laugh out loud! [Stunned]	You sing so beautifully, haha [speechless]	You’re singing really well, haha (an awkward smile).

Table 5 shows both Google and DeepL failed to translate the pragmatic function of “lol” accurately, instead keeping the [no comment] emoji. But the table doesn’t show the actual output format in real interfaces. **Figure 1** is a Google Translate screenshot: the system kept the emoji but also translated it alongside the literal phrase “sang very well.”



Figure 1. The abbreviation “lol” and its corresponding emoji

Both Google and DeepL failed to recognize the awkward or sarcastic tone of “lol” in this context; instead, they translated it literally as “laugh loudly” or retained it as “haha,” while interpreting “sang very well” as sincere praise. Although the emoji “[speechless]” was preserved, its lack of contextual guidance made it difficult for Chinese readers to determine whether it accompanied the praise or expressed helplessness. Consequently, the sentence’s meaning became ambiguous: readers couldn’t tell whether the user was genuinely praising or jokingly mocking. This inappropriate translation strategy directly weakened the pragmatic effect, constituting a typical pragmatic linguistic failure.

4.3. Sociopragmatic Failures

4.3.1. Misprocessing of Cultural Taboos and Sensitive Topics

One day, a hilarious video surfaced on TikTok featuring a young man wearing a wig and vintage attire, mimicking his grandmother’s clumsy daily routines—like struggling to stand up from the sofa only to sit back down, fumbling for his phone while wearing reading glasses only to find it right in his hand, or practicing dance moves in front of a mirror only to end up falling.

Table 6. Translation Comparison of the Death Slang “Kick the Bucket”

addresser	text	Google Translation	DeepL Translation	Manual Translation (Reference)
user	Grandma finally kicked the bucket, RIP lol.	Grandma has finally passed away. Rest in peace, ha ha.	Grandma finally kicked the bucket—RIP ha ha.	The old lady finally kicked off her walk; she laughed heartily and passed away peacefully.

Table 6 shows that Google’s literal “kick the bucket”（踢了水桶）and DeepL’s “RIP”（去世）both miss the satirical edge of the death slang, while the human version “deng li er le”（蹬腿儿了）gets the black humor right. **Figure 2** shows actual Google Translate outputs, where “kicked the bucket” and “RIP ha ha” appear side by side. Chinese readers can’t tell that the original meant to mock comedic characters through absurd humor.



Figure 2. Screenshot of Google Translate’s translation of death slang

Machine translation’s real problem is it can’t grasp the context of humorous videos—it takes exaggerated jokes as real death reports. It misses the pragmatic intent behind black humor, treating crafted contrasts as logic errors. And it lacks cultural adaptability, so it can’t produce naturally consistent Chinese wording. These sociopragmatic failures show that translating humor and cultural taboos can’t be fixed by just matching words one to one.

4.3.2. Deviation in Politeness Strategies and Face-saving

Politeness strategies differ across cultures. Chinese often uses negative politeness—self-deprecation or indirect language to save the other person’s face. English prefers positive politeness, like direct praise or a relaxed tone to build rapport^[7]. Machine translation often copies source-language polite phrases word for word, ignoring target cultural norms, which makes outputs sound stiff or even rude.

Table 7. Comparing Translations of Chinese Self-deprecating Humility

addresser	text	Google Translation	DeepL Translation	Manual Translation
user	The blogger spoke exceptionally well. I’m just a humble follower, offering a few thoughts—please feel free to point out any inaccuracies.	The blogger’s explanation is excellent. I am just a humble person and would like to offer a few words. Please be gentle if I am wrong.	The blogger speaks so well. I am not talented, just saying a few words. Please pat me lightly if I am wrong.	Great post, buddy. I’m no expert, just sharing a couple of thoughts — please be gentle with me if I’m off.

Table 7 shows the issue clearly. The human translation (“Great post, buddy...”) sounds easygoing but “buddy” may be too casual, so it doesn’t fully capture the original’s modest social intent. So we revised it to “This is a great post... go easy on me.” This feels more natural in English social media—it shows appreciation and gently asks for tolerance through self-deprecating humor, fitting English positive politeness better. **Figure 2** shows Google Translate’s actual output: the deliberately humble tone comes off as awkward in English.

5. Discussion

5.1. Negative Impact on Pragmatic Failures

The pragmatic failures in machine translation of cross-border social media comments exert multifaceted negative impacts on cross-cultural communication. First, there’s a gap between what’s said and what’s meant: irony comes off as sincere praise, and self-deprecation looks like real shame, so people misread the true attitude^[8]. Second, these pragmatic failures can easily trigger negative emotions and cultural stereotypes. When cultural taboos are translated literally, target-language users may get offended and blame that rudeness on the source culture. Studies also show that for ethically sensitive content, users rate MT much lower than human translation^[8]. Moreover, frequent failures hurt users’ trust and lower their willingness to do cross-cultural interactions. In the long run, these bad experiences create a loop: “unreliable MT→less cross-cultural talk→deeper stereotypes.”

5.2. Implications

It’s clear that talent competencies need restructuring^[9]. First, translation teaching should go beyond just vocabulary and grammar by adding a pragmatic awareness module. Using real social media cases, students compare human and machine translations, spot typical failures like irony, self-deprecation, internet slang, and cultural taboos, and practice post-editing. For example, they can compare the Google translation in Table 7 with the improved human version and discuss why “go easy on me” works better in English social media than “please be gentle,” helping them see the value of pragmatic equivalence.

Second, building that specific competence means focusing on context and cultural relativity. Using role-playing and case studies, students learn how the same politeness strategy changes across cultures. They should also be trained to spot potential failures and use contextual clues when using MT to lower the risk of misunderstanding. Future research could use acceptance surveys to back up the optimized translations proposed here.

6. Conclusion

This study uses cross-border reviews from TikTok and Xiaohongshu as a corpus to look at pragmatic failures in machine translation. After analyzing 200 reviews with both quantitative and qualitative methods, we found that mainstream neural MT systems often make two types of errors: pragmalinguistic failures—like literal translation of internet slang, loss of irony, and missing abbreviation meanings—and sociopragmatic failures, such as ignoring cultural taboos, politeness strategies, and self-deprecating humor. Google Translate and DeepL show no big difference in how often these failures occur, which suggests pragmatic weaknesses are a common problem across current NMT systems. The main reasons include mismatched training data, lack of pragmatic reasoning, and poor cross-cultural modeling. These failures cause a gap between what’s meant and what’s communicated, hurt user trust, reinforce stereotypes, and block cross-cultural dialogue. A limitation of

this study is that we only used Chinese-English data from these two platforms, so the findings may not apply directly to other language pairs or social media settings. Future work should expand to more languages, build culturally sensitive databases, and set up human-machine collaborative post-editing. Translation teaching should focus on pragmatic awareness and post-editing skills, while platforms should add user feedback features to improve the cross-cultural adaptability in the field of MT.

Disclosure statement

The author declares no conflict of interest.

References

- [1] Weng Y, Cheng X, 2023, Analysis of failure Types in Machine Translation of Literary Texts: A Case Study of Google's English Translation of China Modern and Contemporary Novels, *Journal of Nanjing University of Science and Technology (Social Sciences Edition)*, 36(1): 65-72.
- [2] Zhao Ruoruiqiong (Trieu Nguyen Nhu Quynh), 2024, A Study on Practical failures in Machine Translation from Chinese to Vietnamese for Cross-Border Trade on the China-Vietnam Route, Guangxi Normal University.
- [3] Li Y, Song X, 2025, Beyond Neutrality: Mapping Two Decades of Research on Machine Translation Bias (2005–2024), *SAGE Open*, 15(4):21582440251392700.
- [4] Lin CA, Park S, Xu X, et al., 2024, Exploring Transcultural Communication via Perceived Social Distance and Intergroup Acceptance toward K-Pop and Asian Culture, *Howard Journal of Communications*, 35(2): 200-216.
- [5] Cheng R, 2025, Cross-border interaction on social media and cross-border language governance online: A case study of the influx of “foreign Douyin refugees” onto Xiaohongshu, *Journal of Language Governance*, 2025(1): 196-208, +233-234.
- [6] Wang Y, Ye D, 2025, The Role of Emojis on Social Media in Cross-Cultural Communication: A Case Study of Chinese and the United States University Students, *SHS Web of Conferences*, 220: 02014.
- [7] Al Khatib MA, 2021, (Im)politeness in Intercultural Email Communication between People of Different Cultural Backgrounds: A Case Study of Jordan and the USA, *Journal of Intercultural Communication Research*, 50(4): 409-430.
- [8] Asscher O, Glikson E, 2023, Human evaluations of machine translation in an ethically charged situation, *New Media & Society*, 25(5): 1087-1107.
- [9] Han N, 2026, The Game Between DeepSeek Machine Translation and Human Translation Under the Theory of Functional Parity: A Case Study of Subtitle Translation for *Ne Zha: The Demon Child Comes into the World*, *Overseas English*, 2026(5): 1-5.
- [10] Tang Y, Zun Z, & Zhai C, 2026, The Transformation of the Translation Industry Under the Impact of Artificial Intelligence: A Study on Response Strategies Led by the Translation Association, *Overseas English*, 2026(5): 21-23.

Publisher's note

Bio-Byword Scientific Publishing remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.