



Navigating the Gap: A Lexico-Syntactic Error Analysis of Human and AI-Powered Translation in English-Indonesia ELT Journal

¹Immatul Jannah Alimir, ²Veni Roza, ³Widya Syafitri

1,2,3 English Language Education, UIN Sjech M. Djamil Djambek Bukittinggi, Indonesia

immatulalimir@gmail.com, veniroza@uinbukittinggi.ac.id, widyasyafitri@uinbukittinggi.ac.id

INFO ARTIKEL

Riwayat Artikel:

Diterima: 01-07-2026

Disetujui: 03-07-2026

Kata Kunci:

Kesalahan leksikal

Kesalahan sintaksis

Abstrak ELT

Terjemahan manusia

Terjemahan AI

Keywords:

Lexical errors

Syntactic errors

ELT abstract

Human translation

AI translation

ABSTRAK

Abstrak : Penelitian ini menginvestasikan kesalahan leksiko-sintaksis dalam terjemahan bahasa Indonesia dari abstrak jurnal Pendidikan Bahasa Inggris (ELT) berbahasa Inggris, dengan membandingkan dari hasil yang diproduksi oleh penerjemah manusia dan sistem berbasis AI. Korpus penelitian terdiri dari sepuluh abstrak bahasa Inggris beserta terjemahannya, sehingga total terdapat dua puluh teks yang dianalisis. Terjemahan manusia disusun oleh mahasiswa pascasarjana yang terlatih dalam studi penerjemahan, sementara terjemahan AI dihasilkan secara otomatis menggunakan alat penerjemahan mesin (machine translation). Setiap abstrak diperiksa baris demi baris, dan segmen terjemahan dikodekan berdasarkan kesalahan pilihan leksikal, kesalahan sintaksis, penghilang (omission), dan penambahan (addition). Analisis mengungkapkan enam belas jenis kesalahan yang dominan, dimana terjemahan manusia menunjukkan frekuensi penyimpangan leksikal dan sintaksis yang lebih tinggi, sedangkan terjemahan AI umumnya lebih luwes namun terkadang tidak akurat dalam terminologi dan penyesuaian tingkat klausa. Tidak ditemukan adanya penghilangan atau penambahan diseluruh korpus, yang menunjukkan bahwa penyimpangan terbatas pada pemilihan kata dan konstruksi kalimat, bukan pada kelengkapan konten. Temuan ini menyoroti pentingnya pilihan leksikal dan struktur sintaksis yang presisi dalam penerjemahan akademik, serta menekankan perlunya penyuntingan pasca-terjemahan (post-editing) oleh manusia ketika memanfaatkan penerjemahan berbantuan AI dalam konteks Pendidikan Bahasa Inggris.

Abstract: This study investigates lexico-syntactic errors in Indonesian translations of English-language ELT journal abstracts, comparing outputs produced by human translators and AI-powered systems. The corpus consists of ten English abstracts and their corresponding translations, resulting in twenty texts analyzed in total. Human translations were prepared by graduate students trained in translation studies, while AI translations were generated automatically using neural machine translation tools. Each abstract was examined line by line, and translation segments were coded for lexical choice errors, syntactic errors, omissions, and additions. The analysis revealed sixteen dominant error instances, with human translations exhibiting a higher frequency of lexical and syntactic deviations, whereas AI translations were generally smoother but occasionally inaccurate in terminology and clause-level alignment. No omissions or additions were observed across the corpus, indicating that deviations were confined to word selection and sentence construction rather than content completeness. These findings highlight the importance of precise lexical choices and syntactic structuring in academic translation and emphasize the need for human post-editing when utilizing AI-assisted translation in the context of English Language Teaching.



<https://doi.org/10.31764/telaah.vxiY.ZZZ>



This is an open access article under the **CC-BY-SA** license

A. INTRODUCTION

The rapid evolution of artificial intelligence (AI) in Natural Language Processing (NLP) has fundamentally restructured the epistemological foundation of Translation Studies (TS), precipitating a paradigm shift in which the demarcation between human and machine agency grows increasingly untenable (O'Hagan, 2019). This dissolution of boundaries is not merely technological but deeply ideological, compelling scholars and practitioners alike to reconceptualize translation as a sociotechnical practice rather than an exclusively cognitive or linguistic endeavor. Within the specific pedagogical context of English Language Teaching (ELT) in Indonesia, this technological transformation intersects with the intensifying demands of global academic publishing in ways that are both structurally consequential and institutionally underexamined.

The translation of research abstracts constitutes a particularly high-stakes discursive genre, distinguished by its extreme lexical density, genre-specific rhetorical conventions, and the non-trivial pragmatic function of positioning scholarly work within international epistemic communities. Such characteristics expose the inherent limitations of both human cognitive processing under conditions of specialized domain knowledge and the architectural constraints of Neural Machine Translation (NMT) systems, which despite remarkable advances in fluency and surface-level coherence, continue to demonstrate systematic deficiencies in preserving disciplinary voice, hedging strategies, and inter textual nuance. As Indonesian higher education institutions increasingly operational publication mandates in Scopes-indexed or Sinta-accredited journals as metrics of institutional internationalization, translation competence has emerged as a critical academic literacy, one whose underdevelopment carries measurable consequences for researchers' visibility, credibility, and participation in global knowledge production (Arsyad & Arono, 2016). Recent advances in neural machine translation (NMT), exemplified by tools such as Google Translate, DeepL, and ChatGPT-based translators, have generated substantial interest and debate regarding their potential to replace or assist human translators. Proponents argue that AI tools

offer speed, consistency, and scalability (Koehn, 2020; Moorkens, 2022).

Error analysis (EA) has long been a central methodology in second language acquisition (SLA) and applied linguistics research. EA provides a structured framework for evaluating linguistic quality and pinpointing areas of divergence between source and target texts. Lexico-syntactic errors, which encompass vocabulary selection mistakes and structural misinformation, are particularly consequential in academic translation as they can distort scientific meaning and undermine the clarity of scholarly discourse.

The present study therefore addresses a significant gap in the literature by conducting a systematic lexico-syntactic error analysis of ELT journal article translations produced by both human professionals and AI-powered tools. By comparing error types and frequencies across translation modalities, this research provides empirically grounded insights relevant to translators, ELT practitioners, language educators, and technology developers. The study is guided by the following research questions: (1) What types of lexico-syntactic errors appear in human and AI-powered translations of English-to-Indonesian ELT journal articles? (2) What is the frequency distribution of these errors across both translation modalities? (3) How do the error patterns in human and AI-powered translations differ in terms of lexico-syntactic characteristics?

The theoretical framework of this study draws on three interrelated bodies of knowledge: error analysis theory (Corder, 1967; James, 1998), translation quality assessment (House, 2015), and neural machine translation evaluation (Koehn, 2020). Lexico-syntactic errors are defined in this study as deviations from the target language norms at the levels of word selection, collocation, phrase structure, clause construction, and sentence organization. These categories are critical in academic texts, where precision and coherence are paramount.

The significance of this study is multidimensional. For ELT researchers and journal editors, it offers practical guidelines for evaluating translation quality. For language learners and teachers in Indonesia, it highlights the limitations and affordances of AI translation tools. For translation

technology developers, it provides linguistic evidence about persistent error patterns requiring improvement. Ultimately, this study contributes to the growing discourse on human-AI collaboration in academic translation and language education.

The present study addresses three converging gaps in the literature. First, there is an insufficient body of research specifically examining human versus AI-generated translations in the ELT journal abstract genre within the English-Indonesian language pair. Second, lexico-syntactic errors in Indonesian-English academic translation remain underexplored relative to their pedagogical significance. Third, as AI translation becomes increasingly embedded in academic workflows, there is a pressing need to evaluate its performance empirically and comparatively. In addressing these questions, the study contributes to the intersection of translation studies, applied linguistics, and educational technology—three fields whose convergence is increasingly significant in the contemporary academic landscape.

B. METHOD

Design

This study employed a qualitative descriptive approach combined with comparative error analysis to investigate the lexico-syntactic differences between human and AI-powered translations in English-Indonesian ELT journal texts. The qualitative approach was selected because the study focused on interpreting linguistic phenomena, contextual meanings, translation patterns, and discourse structures rather than measuring numerical variables statistically. Comparative error analysis was used to identify, classify, and interpret translation errors produced by human translators and AI translation systems.

This research prioritized analysis of how lexical selections and syntactic configurations affect semantic correspondence, legibility, and pragmatic efficacy in scholarly translation endeavors. By means of this analytical framework, the investigation sought to produce an exhaustive appraisal of the capacities and constraints inherent in machine-driven translation methodologies when contrasted with conventional human translation practices. The analytical framework is grounded in Dulay, Burt, and

Krashen's (1982) surface strategy taxonomy, operationalized at the lexico-syntactic level and supplemented with Baker's (1992) framework for lexical equivalence. Four primary error types are distinguished: (1) lexical choice errors, (2) syntactic errors, (3) omission, and (4) addition.

Subject or Participant

In this study, the human translation condition comprises Indonesian versions of the English abstracts listed in the table. These translations were produced by graduate students trained in translation studies, ensuring that the choices made reflect informed human judgment regarding lexical, syntactic, and pragmatic aspects of the source texts. For the AI translation condition, the translation data consist of Indonesian translations generated automatically by AI systems, such as Google Translate, as recorded in the uploaded corpus. Both conditions are treated as parallel target-text renditions of the same English abstracts, allowing for a controlled comparison. By analyzing these two sets side by side, the study is able to systematically identify patterns of lexical choice, syntactic structure, omissions, and additions, thereby highlighting the distinctive characteristics of human versus AI translation strategies. This design ensures that any observed differences in translation quality or error patterns can be confidently attributed to the translation agent rather than variations in the source material, providing a rigorous framework for evaluating English-to-Indonesian academic translation performance.

Data and Source of Data

The data for this study comprises ten English-language abstracts from ELT journals, each accompanied by two Indonesian-language translations, one produced by human translators and the other generated by AI systems. This arrangement results in a total of twenty target texts for detailed examination. The selected abstracts encompass a wide range of topics relevant to English Language Teaching, including the communicative approach, the development of critical thinking skills, informal and digital learning practices, online instructional methods, Technology-Enhanced Language Learning (TELL), the use of authentic teaching materials, ICT integration, and the application of artificial intelligence in language education. Given the

relatively small size of the corpus and its purposive sampling, the focus of the study is on analytical depth and a nuanced qualitative evaluation of lexico-syntactic errors, rather than on producing results that can be generalized statistically across the entire ELT literature.

Data Collecting Technique

Each English abstract was compared line by line with the human and AI versions. Segments were then coded for lexical choice error, syntactic error, omission, and addition. A coding mark was assigned only when a deviation was clearly traceable to the source text. In the final coding of this corpus, omission and addition did not appear as dominant categories, so the results focus mainly on lexical and syntactic deviation.

Data Analysis Technique

The coded instances were counted by condition and by error type. Because this is a small qualitative corpus, the counts are treated as descriptive frequencies rather than inferential statistics. The goal is to show patterns of translation behavior, not to claim population-wide estimates.

C. RESULT AND DISCUSSION

Result

A total of 16 dominant lexico-syntactic error instances were identified across the 20 Indonesian translations included in the corpus. Of these, 10 errors (62.5%) were observed in human translations, while 6 errors (37.5%) were found in AI-generated translations. The analysis revealed that lexical choice errors and syntactic errors were the only categories that consistently appeared in the corpus, whereas omission and addition errors were absent in both conditions, reflecting the faithful rendering of all elements from the source abstracts in the final translations.

Table 1. Distribution of Lexico-Syntactic Errors: Human vs. AI Translation

Error Type	Human (n)	Human (%)	AI (n)	AI (%)	Total
Lexical Choice Errors	6	60.0%	3	50.0%	9

Syntactic Errors	4	40.0%	3	50.0%	7
Omission Errors	0	0.0%	0	0.0%	0
Addition Errors	0	0.0%	0	0.0%	0
Total	10	100.0%	6	100.0%	16

The table indicates that lexical choice errors dominated the human translations, constituting 60% of all observed errors. This pattern suggests that human translators were slightly more prone to mis-selection of target-language equivalents, overgeneralization, or minor register mismatches, particularly in discipline-specific terminology. In contrast, AI translations exhibited an even distribution between lexical and syntactic errors, each accounting for 50% of the AI errors. This balanced pattern reflects AI’s tendency to produce grammatically coherent sentences while occasionally failing in precise lexical selection, especially with technical terms or context-sensitive expressions.

Human translators tended to exhibit literal mapping of English clause structures into Indonesian, resulting in minor syntactic infelicities, whereas AI outputs were more balanced syntactically but occasionally suffered from register mismatches or less idiomatic phrase selection. Importantly, the absence of omission and addition errors in both conditions demonstrates that both human and AI translations preserved the complete content of the source abstracts, with deviations largely confined to lexical choice and sentence construction rather than the inclusion or exclusion of material.

Overall, these findings suggest that human translators are slightly more susceptible to lexical infelicities, particularly with discipline-specific terminology, whereas AI translations, while generally more fluent, can be limited by terminology handling and nuanced phrase selection. The pattern observed underscores the importance of human post-editing in AI-assisted translation and highlights the pedagogical value of training translators to refine both lexical and syntactic fidelity in English–Indonesian ELT contexts.

Discussion

In the analysis of the human-translated texts, particularly in ELT-03 and ELT-04, several lexical challenges were observed. These included overly literal translations of abstract concepts and

collocations that did not align naturally with idiomatic academic Indonesian, resulting in phrasing that, while understandable, lacked fluency and elegance. Similarly, in ELT-07, the human translation exhibited a distinct lexical mismatch, where a specialized ELT term was rendered in a way that diverged from its precise disciplinary meaning.

AI-generated translations, on the other hand, generally produced smoother surface-level syntax, yet they encountered difficulties with items that were highly domain-specific. For example, in ELT-06, the acronym TELL was incorrectly interpreted, and in ELT-08, the metaphorical expression home/host was not conveyed with sufficient semantic accuracy, leading to subtle shifts in meaning.

Regarding syntactic fidelity, errors in human translations were mainly characterized by calque-like clause structures and the direct transfer of English word order into Indonesian. Although these constructions remained readable, they occasionally failed to reflect the idiomatic norms of academic Indonesian, particularly in complex or multi-clause sentences. AI syntactic errors were less frequent but still appeared in texts containing denser syntactic arrangements. These issues primarily manifested as overly literal alignments, rather than complete grammatical breakdowns.

Because omission and addition errors were absent in the coded corpus, it can be inferred that the principal gap in translation quality arises not from missing or extraneous content but from lexical precision and syntactic naturalness. This finding is especially salient for ELT abstracts, in which concise wording and well-structured clause relationships are critical to accurately conveying nuanced academic meaning. The results underscore the importance of both careful lexical selection and syntactic structuring in producing translations that are not only accurate but also fluent and contextually appropriate in Indonesian academic discourse.

D. CONCLUSION AND SUGGESTION

This study investigated ten ELT journal abstracts, comparing human-generated and AI-generated Indonesian translations through a systematic lexico-syntactic error analysis. Across the

twenty translated texts, the coding process identified 16 dominant error instances, all of which fell into the categories of lexical choice or syntactic deviations. Overall, human translations exhibited a higher frequency of both lexical and syntactic discrepancies, while AI translations, although often more fluent at the surface level, remained susceptible to terminological inaccuracies and clause-level misalignment.

The findings emphasize the pedagogical implications for translation instruction in the ELT context. Specifically, training programs should place explicit emphasis on the accurate rendering of specialized ELT terminology, appropriate collocations, and effective clause restructuring. Furthermore, students should be encouraged to engage in post-editing of AI-generated translations, rather than treating them as finalized outputs, to ensure that both precision and idiomatic register are maintained. These results highlight that, within academic translation, surface fluency alone is insufficient, and meticulous attention to lexical choice and syntactic organization remains crucial for maintaining the integrity of meaning.

For future research, it is recommended to expand the corpus to include additional journals and a broader array of ELT abstracts, as well as to conduct a more systematic comparison across different AI translation systems. A larger and more diverse data set would allow researchers to assess whether the patterns of lexical and syntactic error observed in this study generalize across various ELT genres and whether AI performance varies consistently in specialized academic contexts. This approach would further inform both translation pedagogy and the development of AI-assisted translation tools tailored to English-Indonesian academic translation.

ACKNOWLEDGMENT

The authors would like to express their gratitude to the English Language Education Department, UIN Sjech M. Djamil Djambek Bukittinggi, and to all parties who contributed to the completion of this research.

REFERENCES

- [1] Ardhito, A., Purwarianti, A., & Mahendra, R. (2020). Neural machine translation for Indonesian: Challenges and

- preliminary results. *Proceedings of the 2020 International Conference on Asian Language Processing*, 45–50.
- [2] Baker, M. (1992). *In other words: A coursebook on translation*. Routledge.
- [3] Bentivogli, L., Bisazza, A., Cettolo, M., & Federico, M. (2016). Neural versus phrase-based machine translation quality: A case study. *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, 257–267.
- [4] Castilho, S., Moorkens, J., Gaspari, F., Calixto, I., Tinsley, J., & Way, A. (2017). Is neural machine translation the new state of the art? *The Prague Bulletin of Mathematical Linguistics*, 108(1), 109–120.
- [5] Corder, S. P. (1967). The significance of learners' errors. *International Review of Applied Linguistics in Language Teaching*, 5(4), 161–170.
- [6] Dardjowidjojo, S. (2003). English language teaching in Indonesia. *EA Journal*, 21(1), 22–30.
- [7] Dulay, H., Burt, M., & Krashen, S. (1982). *Language two*. Oxford University Press.
- [8] Emilia, E. (2008). *Menulis tesis dan disertasi*. Alfabeta.
- [9] House, J. (1997). *Translation quality assessment: A model revisited*. Gunter Narr Verlag.
- [10] Hyland, K. (2000). *Disciplinary discourses: Social interactions in academic writing*. Longman.
- [11] Koehn, P. (2017). *Neural machine translation*. Cambridge University Press.
- [12] Koehn, P., & Knowles, R. (2017). Six challenges for neural machine translation. *Proceedings of the First Workshop on Neural Machine Translation*, 28–39.
- [13] Landis, J. R., & Koch, G. G. (1977). The measurement of observer agreement for categorical data. *Biometrics*, 33(1), 159–174.
- [14] Lauder, A. (2008). The status and function of English in Indonesia: A review of key factors. *MAKARA Sosial Humaniora*, 12(1), 9–20.
- [15] Mackey, A., & Gass, S. M. (2005). *Second language research: Methodology and design*. Lawrence Erlbaum Associates.
- [16] Nord, C. (1997). *Translating as a purposeful activity: Functionalist approaches explained*. St. Jerome Publishing.
- [17] Popovic, M., & Ney, H. (2007). Word error rates: Decomposition over POS classes and applications for error analysis. *Proceedings of the Second Workshop on Statistical Machine Translation*, 48–55.
- [18] Pym, A. (1992). Translation error analysis and the interface with language teaching. In C. Dollerup & A. Loddegaard (Eds.), *Teaching translation and interpreting* (pp. 279–288). John Benjamins.
- [19] Setiawan, B., Wulandari, D., & Permadi, A. (2021). Machine translation accuracy in Indonesian academic texts: A comparative study. *BAHTERA: Jurnal Pendidikan Bahasa dan Sastra*, 20(2), 112–125.
- [20] Shreve, G. M., & Angelone, E. (Eds.). (2010). *Translation and cognition*. American Translators Association.
- [21] Sneddon, J. N., Adelaar, A., Djenar, D. N., & Ewing, M. C. (2010). *Indonesian: A comprehensive grammar* (2nd ed.). Routledge.
- [22] Swales, J. M. (1990). *Genre analysis: English in academic and research settings*. Cambridge University Press.
- [23] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, L., & Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
- [24] Vilar, D., Xu, J., D'Haro, L. F., & Ney, H. (2006). Error analysis of statistical machine translation output. *Proceedings of LREC 2006*, 697–702.
- [25] Wu, Y., Schuster, M., Chen, Z., Le, Q. V., Norouzi, M., Macherey, W., & Dean, J. (2016). Google's neural machine translation system: Bridging the gap between human and machine translation. *arXiv preprint arXiv:1609.08144*.