

Google Translate Versus Matecat for Religious Text Translation: A Study of Iranian Students' Speed, Accuracy, and Perceptions¹

____ Abolfazl Khorasanizadeh Gazki² & Dariush Nejad Ansari Mahabadi³

Abstract

Technological advancement has led to the advent of numerous Machine Translation system and Computed-Assisted Translation (CAT) tools. This study compared the effectiveness of Google Translate as an MT system with Matecat, a CAT tool. It examined their impact on translation quality, speed, and user feedback on both systems. The research involved two classes at the Islamic Azad University of Qom, with 16 students assigned to the Matecat group and 11 to the Google Translate group. All participants first translated a 250-word religious text using dictionaries and completed a placement test showing they shared an intermediate English proficiency level. Following instructions, participants used their assigned system to translate the same text for the post-test. The research team assessed translation quality using Waddington's model. Dependent t-tests showed that while Google Translate significantly reduced translation time without improving quality, Matecat achieved faster and better quality than human translation. Independent t-tests found no significant differences between the systems regarding translation accuracy and speed. Students responded positively to both systems, noting their user-friendly interfaces and accurate religious terminology and grammar handling. They expressed satisfaction with both tools and indicated they would continue using them.

Keywords: Google Translate, Machine Translation, Machine Translation Output, Matecat, Translation Quality Assessment

1. This paper was received on 22.11.2024 and approved on 17.01.2025.

2. Corresponding Author: M.A. in Translation Studies, Department of English Language and Literature, Faculty of Foreign Languages, University of Isfahan, Isfahan, Iran; email: abolfazlkhosani@fgn.ui.ac.ir

3. Associate Professor, Department of English Language and Literature, Faculty of Foreign Languages, University of Isfahan, Isfahan, Iran; email: d.nejadansari@fgn.ui.ac.ir

1. Introduction

Machine Translation (MT) systems and Computer-Assisted Translation (CAT) tools have become increasingly popular among users and researchers (professionals) alike, providing various options and interfaces to cater to different needs. The rise of these systems has prompted a closer examination of their effectiveness, particularly in educational settings. This research specifically compared two prominent systems: Google Translate, a machine translation system, and Matecat, a computer-assisted translation (CAT) tool. The goal was to identify which platforms are more effective and user-friendly for English as a Foreign Language (EFL) students, who often encounter difficulties selecting and utilizing MT and CAT tools due to the overwhelming number of choices available.

EFL students frequently struggle with the decision-making process of choosing the right system. The vast array of options can lead to confusion and frustration, making it essential to evaluate which systems provide the most support and ease of use. This study aimed to shed light on the most beneficial system by considering each platform's technological advancements and the users' perspectives. Understanding user experiences is crucial in determining the overall effectiveness of these tools in educational contexts.

In addition to comparing the systems, the research also sought to assess the quality of the translations produced by these systems after they have been edited. It is essential to evaluate whether the initial output from these systems is valuable enough to justify the time and effort spent on editing. This aspect of the study is particularly relevant in a world where efficiency and accuracy are paramount in translation tasks. By examining the post-editing process, the research aimed to provide insights into the practical applications of these systems in real-world scenarios.

The impact of these two systems on translation speed and quality has been a significant concern since the inception of these technologies. This study explored the

critical factors of timing and quality, delving into how these elements affect user satisfaction and the overall effectiveness of the two systems. By addressing these concerns, the research aimed to contribute to the ongoing discourse surrounding the role of MT and CAT tools in modern translation practices.

The present study aims to answer the following research questions:

1. To what extent can "Google Translate and Matecat" improve the quality and timing of Iranian students' English-to-Persian translation?
 2. Are there any significant differences between the MT systems, "Google Translate and Matecat," regarding the quality and timing of users' translations?
 3. What are the users' attitudes towards using "Google Translate and Matecat"?
2. Literature Review

This section reviews the studies that were conducted that are pertinent to the present study.

The study by Doherty and Kenny (2014) involved designing and evaluating a Statistical Machine Translation (SMT) syllabus for postgraduate translation students. A mixed-methods approach was used to gather data on students' views of the syllabus and their self-assessment of learning. They highlighted that students moved from a superficial understanding of SMT to a deeper appreciation of its complexities.

Purwaningsih (2016) found that Google Translate has beneficial aspects and is particularly useful for individuals seeking rapid translation. Google Translate is designed to assist the reader in deciphering the meaning of the target text.

Laubli et al. (2019) examined the impact of Neural Machine Translation (NMT) on translation speed and quality in the banking and financial sector. NMT allowed translators to work faster. Rather than starting from scratch, they used domain-specific memories and terminology. NMT didn't lower quality. Translators spent more time on stylistic changes, highlighting the need for training when using NMT for productivity.

Macken et al. (2020) found that machine translation (MT) offers quantifiable benefits in practical contexts, with neural MT (NMT) showing more consistent benefits than statistical MT (SMT). This supports the findings that neural systems generally produce more meaningful results. Their research's key advantage and drawback was using genuine translations under typical conditions, leading to accurate time estimates for fewer segments and neglecting other duties like project management and quality assurance.

Tasmedir et al. (2023) found that teachers shared concerns similar to those in previous studies but were more favorable to using machine translation (MT) in instruction. Teachers feared that students using MT for writing tasks without proper analysis would affect their learning. Despite initial negative views, teachers later showed positive attitudes towards MT, likely due to favorable sentiments about technology's role and proactive use of translation-based methods.

According to Pal et al. (2023), by utilizing target variables to model phoneme lengths, the translation quality experiences a significant improvement of 1.8 BLEU, while the speech overlap is enhanced by 0.023 compared to the interleaved baseline.

According to Xu (2024), automated evaluation models significantly enhance translation system performance, achieving correctness rates of 94.8% on Google Translate and 92.6% on Wikipedia datasets. These rates slightly surpass manual evaluation rates, indicating the effectiveness of the proposed methods.

3. Methodology

This section is divided into four subsections: "Design," "Instruments," "Participants," and "Procedure."

3.1. Design

The study employed a pre-experimental research design, incorporating both a pre-test and a post-test. Participants were divided into two experimental groups,

one using the MT system "Google Translate" as an MT system and the other using an actual CAT tool, "Matecat."

3.2. Instruments

The researchers used multiple instruments to collect data: a demographic questionnaire, the Oxford Placement Test to measure language proficiency, and pre-test/post-test translation tasks to assess translation skills. They also employed Google Translate (an MT system) and Matecat (a CAT tool) to evaluate the effectiveness of these technologies in translation tasks. An MT system automatically translates text between languages, while a CAT tool assists human translators with features to improve efficiency and accuracy.

Additionally, Waddington's Model for Translation Quality Assessment was used to evaluate translation quality, and a researcher-made attitude survey gathered insights into participants' perceptions of translation practices. This comprehensive approach enabled a detailed analysis of both quantitative and qualitative aspects of translation performance and attitudes, offering a deeper understanding of factors influencing translation quality.

3.3. Participants

Participants were selected from the Islamic Azad University of Qom by randomly choosing two classes and enrolling their students. The "MateCAT" group included sixteen people, while the "Google Translate" group comprised eleven.

3.4. Procedure

In the study's initial phase, participants translated an English religious text into Persian using only dictionaries, with their translation times recorded. In the next phase, they translated the same text using Google Translate and Matecat. The Google Translate group post-edited the output, while the Matecat group post-edited the system-generated translation. Both processes were timed, and participants completed

a researcher-designed questionnaire to assess their attitudes toward the translation process.

Two independent raters evaluated the translations for accuracy and impartiality to ensure data reliability. SPSS was used to analyze pre-test and post-test results, offering insights into performance and translation method effectiveness.

Before the pre-test, participants took the Oxford Placement Test (60 items: vocabulary, cloze tests, reading comprehension) to establish a language proficiency baseline, which was correlated with their translation performance. This structured approach ensured a comprehensive understanding of participants' capabilities and the impact of machine translation systems and CAT tools on their work.

4. Results and Discussion

This section is divided into two sub-sections, "Results" and "Discussion."

4.1. Results

Here are the results for the first research question.

To investigate the improvement of translation skills in EFL students, the researcher broke down the question into translation quality and timing. A Kappa Cohen test was used to assess inter-rater reliability for the translation quality of Google Translate users, as two raters evaluated the scores. Pre-test and post-test scores were analyzed for inter-rater reliability with a Kappa Cohen test.

Table 4.1 showed the Kappa Cohen test results for the Google Translate group's Pre-test. The hypothesis of no significant agreement between the two raters was rejected, indicating they agreed. The rejection was supported by a P-value of less than 0.05, precisely 0.000. The Kappa value of 0.383 correlated with the raters' agreement.

Table 4.1. Kappa Cohen (inter-rater reliability) test from the pre-test translation task done by Google Translate

		Value	Asymptotic Standard Error, ^a	Approximate T ^b	Approximate Significance
Measure of Agreement	Kappa	0.383	0.155	3.949	0.000
N of Valid Cases		11			
a. Not assuming the Null hypothesis.					
b. Using the asymptotic standard error, assuming the Null hypothesis					

Table 4.2 showed the Kappa Cohen test results for the Google Translate group's Post-test. The hypothesis of no significant agreement between the raters was rejected, indicating their agreement, supported by a P-value of 0.000. The Kappa value of 0.161 showed a direct relationship in their agreement. The mean scores from the two raters were used for the tests.

Table 4.2. Kappa Cohen (inter-rater reliability) test from the post-test translation task done by Google Translate

		Value	Asymptotic Standard Error	Approximate T ^b	Approximate Significance
Measure of Agreement	of Kappa	0.161	0.103	3.602	0.000
N of Valid Cases		11			
a. Not assuming the Null hypothesis.					
b. Using the asymptotic standard error assuming the Null hypothesis.					

The t-test p-value of 0.524 indicated no significant difference between the pre-test and post-test scores of the Google Translate group, implying no improvement in translation quality. Therefore, Google Translate needs to enhance its accuracy in translating religious materials. The next step is to assess its impact on timing, requiring a normality test before conducting a paired sample T-test.

Table 4.3. Paired sample t-test from the pre-test and post-test translation task done by Google Translate

Paired Differences						t	df	Sig. (2-tailed)
Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference					
			Lower	Upper				
Mean of the pre-test—Mean of the post-test	-4.90909	24.68787	7.44367	-21.49463	11.67645	-0.659	10	0.524

Table 4.4, the p-value of 0.004, less than 0.05, indicated a significant difference in the timing of the pre-test and post-test tasks. Both the Lower and Upper Intervals of the Difference were positive, showing that mean timing in the pre-test was longer than in the post-test. This confirms that Google Translate successfully reduced the duration of translation tasks, indicating an improvement.

Table 4.4. Paired sample t-test from the pre-test and post-test timing of the translation task done by Google Translate

	Paired Differences					t	df	Sig. (2- tailed)
	Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference				
				Lower	Upper			
The timing of the participant in translating the pre-test - The timing of the participant in translating the post-test	7.67909	6.94702	2.09461	3.01202	12.34616	3.666	10	0.004

Here are the results for the second research question.

The second research question is divided into quality and timing, similar to the Google Translate group. The Matecat group also uses two raters to assess the quality

of pre-test and post-test tasks. The first step is to conduct a Kappa Cohen test to evaluate inter-rater reliability for both the Pre-test and Post-test scores, beginning with the Pre-test.

Table 4.5 showed the Kappa Cohen test results for the Matecat group's Pre-test. The hypothesis of no significant agreement between the raters was rejected, indicating they agreed. This conclusion was supported by a p-value of 0.000 and a Kappa value of 0.257, showing a clear positive correlation in their agreement.

Table 4.5. Kappa Cohen (inter-rater reliability) test from the pre-test translation quality scores obtained from Matecat

		Value	Asymptotic Standard Error	Approximate T ^b	Approximate Significance
Measure of Kappa Agreement		0.257	0.105	4.126	0.000
N of Valid Cases		16			
a. Not assuming the Null hypothesis.					
b. Using the asymptotic standard error assuming the Null hypothesis.					

Table 4.6 presented the Kappa Cohen test outcomes for the Matecat group's post-test task. The hypothesis of no significant agreement between the raters was rejected, indicating they agreed. This was supported by a p-value of 0.009. However, the Kappa value of 0.097 showed a clear and direct correlation in their agreement.

Table 4.6. Kappa Cohen (inter-rater reliability) test from the post-test translation quality scores obtained from Matecat

		Value	Asymptotic Standard Error	Approximate T ^b	Approximate Significance
Measure of Kappa Agreement		0.097	0.059	2.619	0.009
N of Valid Cases		16			
a. Not assuming the Null hypothesis.					
b. Using the asymptotic standard error assuming the Null hypothesis.					

Table 4.7 showed the t-test results for the Matecat group's pre-test and post-test translation scores. The p-value of 0.023, less than 0.05, indicated a significant disparity in quality. The Lower and Upper Intervals of the Difference were negative, suggesting the mean pre-test scores were lower than the post-test scores. This confirms that Matecat effectively improved translation quality, showing progress.

Table 4.7. Paired sample t-test from the pre-test and post-test of the translation task done by Matecat

		Paired Differences							
		Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference		t	df	Sig. (2-tailed)
					Lower	Upper			
Mean of pre-test - Mean of post-test	-13.15625	20.79160	5.19790	-24.23531	-2.07719	-2.531	15	0.023	

Table 4.8 showed the t-test results for the Matecat group's pre-test and post-test timing data. The p-value of 0.001, less than 0.05, indicated a statistically significant disparity in timing. Both the Lower and Upper Intervals of the Difference were positive, meaning the mean timing for the pre-test was longer than the post-test. Therefore, Matecat effectively reduced the time required for translation tasks, showing improvement.

Table 4.8. Paired sample t-test from the pre-test and post-test of the timing of the translation task done by Matecat

Paired Differences							
Mean	Std. Deviation	Std. Error Mean	95% Confidence Interval of the Difference		t	df	Sig. (2-tailed)
			Lower	Upper			

The timing of the participant in translating the pre-test - The timing of the participant in translating the post-test	11.90875	11.31096	2.82774	5.88156	17.93594	4.211	15	0.001
--	----------	----------	---------	---------	----------	-------	----	-------

Here are the results for the third research question.

The third research question was divided into quality and timing components, followed by a comparison of the two machine translation systems. A paired sample t-test compared the pre-test quality for both Google Translate and Matecat groups, and the same analysis was conducted for the post-test.

Table 4.9 showed the results of an independent paired sample t-test on the Post-test translation scores of both groups, with mean ratings from two raters. The p-value of the Levene test (0.985) exceeded 0.05, indicating that variances were equal. The "t" test was insignificant, with a p-value of 0.887, also exceeding 0.05. This upheld the Null hypothesis, indicating no significant disparity between the two groups' post-test translation performance when using the specific Machine Translation (MT) system.

Table 4.9. Independent sample t-test from the post-test of the translation task done by Google Translate and post-test of the translation task done by Matecat

		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	t	df	Sig. (2- tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
MeanPosttest GTandM	Equal variances assumed	.000	.985	.144	20	.887	1.18182	8.22363	-15.97237	18.33601
	Equal variances are not assumed.			.144	19.740	.887	1.18182	8.22363	-15.98685	18.35049

Table 4.10. Independent sample t-test from the post-test of the timing of the translation task done by Google Translate and post-test of the timing of the translation task done by Matecat

		Levene's Test for Equality of Variances		t-test for Equality of Means						
		F	Sig.	t	df	Sig. (2-tailed)	Mean Difference	Std. Error Difference	95% Confidence Interval of the Difference	
									Lower	Upper
The timing of the participant in translating the Post-test	Equal variances assumed	.316	.580	-1.271	20	.218	-3.04000	2.39184	-8.02929	1.94929
	Equal variances are not assumed.			-1.271	19.704	.219	-3.04000	2.39184	-8.03411	1.95411

Table 4.10 showed the results of an independent paired sample t-test on the post-test translation timing for both groups. The Levene test's p-value was 0.580, exceeding 0.05, indicating that the equality of variances was accepted. The "t" test was also insignificant, with a p-value of 0.218, exceeding 0.05. This upheld the Null hypothesis, indicating no significant disparity in post-test translation timing between the two groups.

Here are the results for the fourth research question.

Here, the reliability of the questionnaire given to the Google Translate group is computed using Cronbach's Alpha.

The Cronbach's Alpha coefficient of 0.844 indicates a high level of reliability, surpassing the satisfactory threshold of 0.7. This suggests that the questionnaire had an adequate level of reliability. Here are the response frequencies from the Google Translate group's attitude questionnaire.

Table 4.11. The Alpha Cronbach's Reliability test results for the Google Translate group's attitude questionnaire

Cronbach's Alpha	N of Items
0.844	25

Table 4.12 shows the frequencies of the answers the Google Translate group's participants provided for the attitude questionnaire.

Table 4.12. Questions of the Google Translate attitude questionnaire

Questions	Completely Agree	Agree	No Opinion	Disagree	Completely Disagree
1. Google Translate can be helpful for quickly producing a product.	27.3%	63.6%	9.1%	0%	0%
2. This course on how to use Google Translate and how to post-edit the output positively influenced my	9.1%	45.5%	18.2%	9.1%	18.2%

attitude toward using MT systems.					
3. The use of MT systems improved my quality of translation.	27.3%	54.5%	18.2%	0%	0%
4. I will use the MT system more often after this course.	9.1%	27.3%	36.4%	9.1%	18.2%
5. I learned how to use the MT system correctly after this course.	27.3%	36.4%	27.3%	0%	9.1%
6. MT system use is a barrier to my creativity.	18.2%	9.1%	18.2%	45.5%	9.1%
7. The errors made by MT causes fatigue.	9.1%	36.4%	9.1%	36.4%	9.1%
8. The errors made by MT could be more apparent.	18.2%	27.3%	36.4%	9.1%	9.1%
9. The use of MT hinders my ability to learn how to translate.	27.3%	18.2%	9.1%	18.2%	27.3%
10. I will conduct more research on MT after this course.	27.3%	27.3%	36.4%	0%	9.1%
11. After this course, I will use MT with complete trust.	27.3%	9.1%	27.3%	18.2%	18.2%
12. I would like to learn more about MT after this course.	27.3%	0%	27.3%	27.3%	18.2%
13. The use of MT requires professional learning.	27.3%	36.4%	18.2%	9.1%	9.1%
14. Using the presented MT system (Google Translate) in this course is easy.	27.3%	45.5%	27.3%	0%	0%
15. The user interface of the presented MT system is straightforward to work with.	9.1%	27.3%	36.4%	9.1%	18.2%
16. The MT system allows me to post a text efficiently.	27.3%	36.4%	18.2%	9.1%	9.1%

17. The presented MT system provides the opportunity to arrive at a final translation faster.	18.2%	36.4%	27.3%	9.1%	9.1%
18. The output of the presented MT system is fluent.	9.1%	18.2%	27.3%	18.2%	27.3%
19. The output of the presented MT system can be used without post-editing.	0%	27.3%	18.2%	18.2%	36.4%
20. One should be cautious while using the presented MT system.	36.4%	18.2%	9.1%	0%	36.4%
21. This MT system should not be used.	18.2%	9.1%	27.3%	27.3%	18.2%
22. The MT system is capable of replacing translators in the future.	18.2%	27.3%	36.4%	0%	18.2%
23. This MT system performs well in choosing the correct religious terminology.	18.2%	9.1%	9.1%	18.2%	45.5%
24. This MT system provides a well-structured sentence based on syntax.	0%	18.2%	9.1%	45.5%	27.3%
25. This MT system provides a well-structured sentence based on syntactic relations	0%	36.4%	27.3%	27.3%	9.1%

Here, the reliability of the questionnaire is computed using Cronbach's Alpha.

Here, Cronbach's Alpha computes the reliability of the questionnaire given to the Matecat group.

The Cronbach's Alpha coefficient of 0.755 indicates a high level of reliability, surpassing the satisfactory threshold of 0.7. This suggests that the questionnaire had an adequate level of reliability. Here are the response frequencies from the Matecat group's attitude questionnaire.

Table 4.13. The Alpha Cronbach's Reliability test results for the Google Translate group's attitude questionnaire

Cronbach's Alpha	N of Items
0.755	25

Table 4.14 shows the frequencies of the answers the Matecat group's participants provided for the attitude questionnaire.

Table 4.14. Questions of the Matecat attitude questionnaire

Questions	Completely Agree	Agree	No Opinion	Disagree	Completely Disagree
1. Matecat can help produce a product quickly.	25%	68.8%	0%	6.3%	0%
2. This course on how to use Matecat and post-edit the output positively influenced my attitude toward using MT systems.	12.5%	43.8%	37.5%	6.3%	0%
3. The use of MT systems improved my quality of translation.	6.3%	43.8%	25%	25%	0%
4. I will use the MT system more often after this course.	31.3%	18.8%	37.5%	12.5%	0%
5. I learned how to use the MT system correctly after this course.	37.5%	43.8%	18.8%	0%	0%
6. MT system use is a barrier to my creativity.	6.3%	25%	18.8%	50%	0%
7. The errors made by MT causes fatigue.	0%	31.3%	37.5%	31.8%	0%
8. The errors made by MT are confusing.	6.3%	37.5%	43.8%	12.5%	0%
9. The use of MT hinders my ability to learn how to translate.	0%	43.8%	18.8%	31.3%	6.3%
10. I will conduct more research on MT after this course.	12.5%	56.3%	18.8%	6.3%	6.3%

11. After this course, I will use MT with complete trust.	18.8%	25%	31.3%	18.8%	6.3%
12. I would like to learn more about MT after this course.	12.5%	56.3%	18.8%	6.3%	6.3%
13. The use of MT requires professional learning.	6.3%	37.5%	25%	31.3%	0%
14. Using the presented MT system (Matecat) in this course is easy.	18.8%	68.8%	12.5%	0%	0%
15. The user interface of the presented MT system is straightforward to work with.	31.3%	37.5%	25%	6.3%	0%
16. The MT system allows me to post a text efficiently.	12.5%	43.8%	25%	18.8%	0%
17. The presented MT system provides the opportunity to arrive at a final translation faster.	6.3%	43.8%	31.3%	12.5%	6.3%
18. The output of the presented MT system is fluent.	6.3%	12.5%	31.3%	43.8%	6.3%
19. The output of the presented MT system can be used without post-editing.	0%	18.8%	6.3%	50%	25%
20. One should be cautious while using the presented MT system.	31.3%	50%	18.8%	0%	0%
21. This MT system should not be used.	6.3%	18.8%	62.5%	12.5%	0%
22. The MT system is capable of replacing translators in the future.	12.5%	18.8%	18.8%	25%	25%
23. This MT system performs well in choosing the correct religious terminology.	0%	31.3%	31.3%	31.3%	6.3%
24. This MT system provides a well-structured sentence based on syntax.	6.3%	25%	18.8%	50%	0%

25. This MT system provides a well-structured sentence based on syntactic relations	6.3%	12.5%	31.3%	31.3%	18.8%
---	------	-------	-------	-------	-------

4.2. Discussion

Macken et al. (2020) observed minimal effects on overall processes but highlighted quality and translation speed enhancements. Their findings align partially with this study's conclusions, which suggest that while Google Translate an MT system does not enhance quality, it can improve timing. Conversely, Matecat, a CAT tool, is capable of boosting both quality and timing. Similarly, Purwaningsih (2016) discovered that Google Translate aids comprehension but fails to deliver high-quality translations. His findings also resonate with this study's results, indicating that Google Translate does not enhance quality, while Matecat does. Laubli et al. (2019) also reported that neural machine translation post-editing saved time and improved quality. This contrasts with this study's assertion that Google Translate cannot enhance quality but can improve timing, while Matecat excels in both aspects.

Tasmedir et al. (2023) explored educators' views on machine translation (MT) in teaching, noting initial concerns about academic integrity, passive learning, and learner autonomy. Exposure to an interactive MT app shifted attitudes positively, highlighting benefits like visualizing language structures and aiding EAL students, though concerns about poor learning habits remained. Unlike Tasmedir's initial skepticism, the current research found more positive outcomes, especially after professional development on effective MT integration.

5. Conclusion

This study evaluated the effectiveness of Google Translate an MT system and Matecat, a CAT tool, in enhancing Iranian students' translation skills from their second to their first language. It aimed to analyze these systems' impact on translation quality, time efficiency, and user attitudes. The results revealed that Google Translate reduced

translation timing but did not significantly improve translation quality. In contrast, Matecat enhanced translation quality and time efficiency by offering advanced features like translation memory and glossaries.

Comparing the overall performance, Google Translate and Matecat were equally efficient in speed. However, Matecat provided superior translation quality for individual users. The study also explored participants' subjective experiences, finding that while both tools were user-friendly, concerns about grammatical accuracy and contextual appropriateness persisted.

The study highlighted the need for critical engagement with MT tools and the importance of proper training. It emphasized that while MT tools offer practical benefits, they may still fall short of human translators' nuanced capabilities. Educators are encouraged to equip students with the skills to use MT tools effectively and critically evaluate their outputs.

6. Works Cited:

- Doherty, S., & Kenny, D. (2014). The design and evaluation of a statistical machine translation syllabus for translation students. *The Interpreter and Translator Trainer*, 8(2), 295–315.
- Laubli, S., Amrhein, C., Düggelein, P., Gonzalez, B., Zwahlen, A., & Volk, M. (2019). Post-editing productivity with neural machine translation: An empirical assessment of speed and quality in the banking and finance domain. In *Proceedings of Machine Translation Summit XVII: Research Track* (pp. 267–272). European Association for Machine Translation.
- Macken, L., Prou, D., & Tezcan, A. (2020). Quantifying the effect of machine translation in a high-quality human translation production process. *Informatics* 7 (2), 1–19.
- Pal, P., Virkar, Y., Mathur, P., Chronopoulou, A., & Federico, M. (2023). Improving isochronous machine translation with target factors and auxiliary counters. *Interspeech 2023*, 37–41. <https://doi.org/10.21437/Interspeech.2023-1063>
- Purwaningsih, D. (2016). Comparing translation produced by Google Translate tool to translation produced by translator. *Journal of English Language Studies*, 1(1).
- Tasdemir, S., Lopez, E., Satar, M., & Riches, N. (2023). Teachers' perceptions of machine translation as a pedagogical tool. *The JALT CALL Journal*, 19(1), 92–112.
- Xu, H. (2024). Assessment of computer-aided translation quality based on large-scale corpora. *International Journal of e-Collaboration (IJeC)*, 20(1), 1–14.