



ISSN: 2617-6548

URL: www.ijirss.com



Improving post-editing of Kazakh translations with fine-tuned large language models: Dataset and evaluation

 Diana Rakhimova¹, Aliya Zhiger^{1,2,3*},  Madina Mansurova¹, Valentin Malykh³, XMagzhan Kairanbay⁴

¹Al-Farabi Kazakh National University, Almaty, Kazakhstan.

²Narxoz University, Almaty, Kazakhstan.

³International University of Information Technologies, Almaty, Kazakhstan.

⁴Til Bilim Institute, Research and Development Department, Qalan, Almaty, Kazakhstan.

Corresponding author: Aliya Zhiger (Email: aliya.zhunussova.zh@narxoz.kz)

Abstract

Machine translation for low-resource languages like Kazakh faces significant challenges due to limited training data, complex morphology, and cultural-linguistic nuances. This paper presents the first comprehensive study on fine-tuning large language models for automated post-editing of Kazakh translations. We introduce KazPE, a systematically annotated dataset containing 10,010 training sentences and 315 test sentences across six domains (medical, scientific, journalistic, oral, fiction, and legal) with detailed error categorization covering 11 linguistic dimensions. Our approach fine-tunes GPT-4.1-mini using supervised learning to improve translation quality through targeted error correction. Human evaluation demonstrates that our fine-tuned model achieves a mean quality score of 0.84 compared to 0.80 for the baseline, representing a 4% relative improvement. The most significant gains occur in morphological-lexical error handling and domain-specific contexts, with legal and medical texts showing improvements of +2.8% and +1.6% respectively. Error analysis reveals that fine-tuning effectively addresses Kazakh's agglutinative morphology and specialized terminology while maintaining performance on error-free sentences. This work establishes the first systematic evaluation framework for Kazakh translation post-editing, providing valuable insights for improving machine translation systems for morphologically rich, low-resource languages. Our dataset, models, and evaluation framework are made publicly available to support future research in Turkic language processing.

Keywords: Fine-tuning, Kazakh, Large language models, Low-resource languages, Machine translation, Morphologically rich languages, NLP, Post-editing.

DOI: 10.53894/ijirss.v8i8.10583

Funding: This research was funded by the Ministry of Science and Higher Education of the Republic of Kazakhstan, grant number BR24993001, “Creation of a large language model (LLM) to maintain the implementation of Kazakh language and increase the technological progress”.

History: **Received:** 6 August 2025 / **Revised:** 9 September 2025 / **Accepted:** 11 September 2025 / **Published:** 9 October 2025

Copyright: © 2025 by the authors. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Competing Interests: The authors declare that they have no competing interests.

Authors' Contributions: All authors contributed equally to the conception and design of the study. All authors have read and agreed to the published version of the manuscript.

Transparency: The authors confirm that the manuscript is an honest, accurate, and transparent account of the study; that no vital features of the study have been omitted; and that any discrepancies from the study as planned have been explained. This study followed all ethical practices during writing.

Publisher: Innovative Research Publishing

1. Introduction

Kazakh, spoken by approximately 13 million people worldwide [1], is a morphologically rich Turkic language belonging to the Kipchak branch of the Turkic language family. Despite its significant speaker population and official status in Kazakhstan, Kazakh remains severely under-resourced in terms of natural language processing tools and datasets compared to high-resource languages such as English, Chinese, or Spanish [2]. This resource scarcity presents unique challenges for machine translation (MT) systems, particularly in handling the language's complex agglutinative morphology, flexible word order, extensive case system, and rich derivational processes [3].

The linguistic complexity of Kazakh poses substantial difficulties for automated translation systems. As an agglutinative language, Kazakh words are formed by adding multiple suffixes to root morphemes, creating potentially unlimited vocabulary through morphological processes. For example, a single Kazakh word can encode information that requires multiple words in English, such as *балаларымыздықи* (balalarymyzdyki) meaning “belonging to our children.” This morphological richness, combined with relatively free word order and a complex system of vowel harmony, creates significant challenges for traditional statistical and neural machine translation approaches [4].

Recent advances in large language models (LLMs) such as GPT-3 [5], GPT-4 [6], and PaLM [7] have demonstrated remarkable zero-shot and few-shot translation capabilities across numerous language pairs. These models have shown particular promise for low-resource languages by leveraging cross-lingual transfer learning and multilingual training objectives [8]. However, despite these advances, LLMs still exhibit significant performance gaps when translating to and from low-resource languages, often producing outputs with semantic inconsistencies, morphological errors, unnatural phrasing, and cultural inappropriateness [9]. The challenges are particularly pronounced for morphologically complex languages like Kazakh, where subtle morphological variations can completely alter meaning.

Post-editing, the process of improving machine-generated translations through human or automated correction, has emerged as a practical solution to bridge the quality gap in MT systems [10]. Unlike traditional approaches that focus on improving the underlying translation model, post-editing works with existing translation outputs to systematically identify and correct errors. This approach is particularly valuable for low-resource languages where training high-quality translation models from scratch is prohibitively expensive due to data scarcity. Traditional post-editing approaches relied heavily on rule-based systems and statistical models [11]. Still, recent work has increasingly explored the use of neural networks and transformer-based models for automated post-editing [12, 13]. The application of fine-tuning techniques to large language models for post-editing represents a promising direction that combines the broad multilingual capabilities of pre-trained models with task-specific adaptation. Fine-tuning allows models to learn specific error patterns and correction strategies while maintaining their general linguistic knowledge [14]. This approach is particularly relevant for languages like Kazakh, where the complex interplay between morphology, syntax, and semantics requires a nuanced understanding that can benefit from both pre-trained multilingual representations and targeted training on language-specific phenomena.

Current research in Kazakh natural language processing has primarily focused on basic tools such as morphological analyzers [3], speech synthesis systems [15], and preliminary machine translation experiments [16]. However, there has been limited work on systematic evaluation of translation quality or development of post-editing systems. The lack of standardized datasets with detailed error annotations has been a significant obstacle to advancing Kazakh Natural Language Processing (NLP) research and comparing different approaches.

In this work, we propose a comprehensive framework for improving Kazakh machine translation through fine-tuned post-editing using GPT-4.1-mini. Our approach addresses several critical gaps in current research by developing both the necessary datasets and methodologies for systematic improvement of Kazakh translation quality. Our work makes the following key contributions:

Dataset Development: We construct and release KazPE, the first large-scale, systematically annotated English-

Kazakh translation dataset with detailed error categorization. The dataset contains 10,010 training sentences and 315 test sentences spanning six distinct domains (medical, scientific, journalistic, oral, fiction, and legal) and covers 11 comprehensive error categories, including semantics, lexical choice, morphology, terminology, style, word order, grammar, calque usage, orthography, idiomatic expressions, and register appropriateness.

Methodological Innovation: We develop a fine-tuning methodology that leverages supervised learning to improve post-editing quality specifically for Kazakh. Our approach uses a conversational prompt-response structure that explicitly guides the model to perform targeted corrections while maintaining semantic fidelity to the source text.

Comprehensive Evaluation Framework: We conduct the first systematic evaluation of LLM performance on Kazakh translation post-editing using human assessment and detailed error analysis across multiple linguistic dimensions. Our evaluation methodology considers both overall quality improvements and specific error category performance.

Empirical Analysis: We provide the first systematic analysis of GPT-4.1-mini's performance on Kazakh translation tasks, identifying key strengths and limitations. Our analysis reveals specific patterns in how the model handles different types of linguistic phenomena in Kazakh.

Practical Impact: Our work establishes baseline performance metrics and provides actionable insights for developing practical translation post-editing systems for Kazakh and related Turkic languages.

Our experimental results demonstrate that fine-tuning GPT-4.1-mini on our annotated dataset yields significant improvements in translation quality, achieving 84% accuracy compared to 80% for the base model. These improvements are consistent across different text domains and complexity levels, suggesting the robustness of our approach. The most dramatic improvements occur in handling morphological-lexical interactions, where our fine-tuned model achieves perfect performance compared to 70% for the baseline, demonstrating the effectiveness of targeted training on Kazakh's complex agglutinative morphology.

The consistency of improvements across diverse domains—with particularly notable gains in specialized fields such as legal and medical texts—underscores both the practical applicability of our approach and its potential for real-world deployment. Our findings provide valuable insights for the broader community working on low-resource and morphologically complex languages, establishing both the potential and the limitations of current fine-tuning approaches for improving translation quality in under-resourced linguistic contexts.

2. Related Works

2.1. Machine Translation for Low-Resource Languages

The challenge of developing effective machine translation systems for low-resource languages has been a central concern in computational linguistics. Low-resource languages, defined as those with limited parallel training data, present unique difficulties that traditional statistical and neural machine translation approaches struggle to address effectively. Early work by Philipp and Rebecca [17] demonstrated that neural machine translation systems require substantially larger datasets than statistical systems to achieve comparable performance, creating particular challenges for languages with limited digital resources.

Recent advances in multilingual machine translation have shown promise for low-resource languages through cross-lingual transfer learning. The work of Johnson, et al. [18] on Google's multilingual neural machine translation system demonstrated that training a single model on multiple language pairs can improve performance for low-resource languages through positive transfer from high-resource pairs. Similarly, the mT5 model by Linting, et al. [19] and the multilingual capabilities of models like mBERT [20] have provided foundations for cross-lingual understanding that benefit low-resource language processing. However, these advances have primarily focused on languages with substantial online presence or those that are closely related to high-resource languages. Morphologically complex languages like those in the Turkic family continue to present significant challenges due to their agglutinative nature, extensive case systems, and complex morphophonological processes that are not well captured by subword tokenization schemes commonly used in neural systems [21].

2.2. Turkic Language Processing and Kazakh NLP

Research on Turkic language processing has highlighted the unique challenges posed by the morphological complexity of this language family. Washington, et al. [3] developed finite-state morphological transducers for Kazakh, providing foundational tools for computational analysis of the language's agglutinative structure. Their work demonstrated the complexity of Kazakh morphology, with extensive case marking, possessive constructions, and derivational processes that create significant challenges for automated processing.

Subsequent work by Altenbek and Wang [4] focused on developing segmentation systems for Kazakh inflective affixes, addressing one of the fundamental preprocessing challenges for computational work with the language. Makazhanov, et al. [15] extended this work by developing open-source tools for Kazakh speech synthesis, contributing to the broader ecosystem of Kazakh language technology. Early machine translation work for Kazakh was primarily rule-based or relied on limited statistical approaches. Badrinath, et al. [22] presented one of the first statistical machine translation systems for Kazakh, though their work was limited by the availability of parallel training data. More recent work by Makhambetov, et al. [16] explored neural machine translation approaches for English-Kazakh translation, showing improvements over earlier methods but still highlighting significant quality gaps compared to high-resource language pairs.

Kartbayev [23] conducted systematic analysis of neural sequence-to-sequence architectures for agglutinative languages including Kazakh, identifying specific challenges related to handling the complex morphological processes that characterize these languages. This work provided important insights into the limitations of standard neural architectures when applied to

morphologically rich languages.

Despite these contributions, Kazakh NLP research has been hampered by the lack of large-scale, systematically annotated datasets. Most existing resources are relatively small and focused on specific tasks, limiting the development of comprehensive evaluation frameworks and comparative studies across different approaches

2.3. Automatic Post-Editing

Automatic post-editing (APE) has emerged as a practical approach to improving machine translation quality without requiring modifications to underlying translation systems. The field originated with rule-based approaches that applied hand-crafted correction patterns to address systematic errors in machine translation output. Early work by Simard, et al. [11] demonstrated that statistical phrase-based methods could be effectively applied to post-editing tasks, treating the problem as a monolingual translation task from "bad" to "good" translations.

The advent of neural approaches to APE marked a significant advancement in the field. Junczys-Dowmunt and Grundkiewicz [12] explored various neural sequence-to-sequence architectures for automatic post-editing, showing that encoder-decoder models could effectively learn to correct systematic translation errors. Their work established important baselines and demonstrated the potential of neural approaches for APE tasks.

More recent work has focused on incorporating additional context and leveraging transformer-based architectures. Correia, et al. [24] developed multi-task learning approaches that combine automatic post-editing with related tasks to improve overall performance. Yang, et al. [25] explored the integration of language models into post-editing systems, showing that pre-trained representations can provide valuable improvements for error correction tasks.

The Workshop on Machine Translation (WMT) has provided important venues for systematic evaluation of APE approaches through shared tasks Chatterjee, et al. [26]. The findings from WMT shared tasks have consistently shown that even modest improvements in automatic post-editing can be valuable, particularly when dealing with domain-specific content or languages where high-quality machine translation systems are not available.

However, most APE research has focused on high-resource language pairs, particularly English-German and English-Spanish. There has been limited systematic investigation of APE approaches for low-resource or morphologically complex languages, creating a significant gap in our understanding of how these techniques transfer to languages with different linguistic characteristics [27].

2.4. Large Language Models for Translation

The emergence of large language models has transformed the landscape of machine translation research. Brown, et al. [5] demonstrated that GPT-3 could perform zero-shot and few-shot translation across numerous language pairs without explicit training on parallel data. This capability emerged from the model's extensive multilingual training and its ability to perform in-context learning from prompt examples.

Subsequent work has systematically evaluated the translation capabilities of large language models. Hendy, et al. [28] conducted comprehensive evaluations of GPT models for machine translation, finding that while these models show impressive zero-shot capabilities, they still lag behind specialized translation systems for high-resource language pairs. However, for low-resource languages, the gap is often much smaller, making LLMs potentially competitive with traditional approaches.

Jiao, et al. [29] specifically examined ChatGPT's translation capabilities, finding that GPT-4 demonstrates significantly improved translation quality compared to earlier versions, particularly for low-resource languages. Their work highlighted the potential of instruction-tuned models for translation tasks and the importance of appropriate prompting strategies.

Research on improving LLM translation performance has explored various fine-tuning approaches. Li and Liang [30] investigated multilingual machine translation with large language models, showing that targeted fine-tuning can substantially improve performance on specific language pairs. Yang, et al. [25] explored the use of large language models for multilingual translation, demonstrating that these models can effectively handle multiple languages simultaneously. Vilar, et al. [31] conducted a detailed analysis of prompting strategies for translation with PaLM, showing that careful prompt design can significantly impact translation quality. Their work provided important insights into how to effectively use large language models for translation tasks and highlighted the importance of evaluation methodologies that account for the unique characteristics of LLM-generated translations.

2.5. Fine-Tuning Approaches for Translation

The application of fine-tuning techniques to improve translation quality has become an active area of research. Traditional fine-tuning approaches for neural machine translation focused on domain adaptation, where models trained on general parallel data were adapted to specific domains or text types [32]. However, the advent of large pre-trained language models has opened new possibilities for fine-tuning approaches.

Recent work by Haoran, et al. [14] proposed paradigm shifts in fine-tuning large language models for machine translation, demonstrating that parameter-efficient fine-tuning methods can achieve substantial improvements while requiring minimal computational resources. Their approach showed particular promise for low-resource language pairs where traditional training approaches are not feasible.

Parameter-efficient fine-tuning methods such as Low-Rank Adaptation (LoRA) by Hu, et al. [33] and prefix-tuning

by Li and Liang [30] have shown promise for adapting large language models to specific tasks without the computational overhead of full fine-tuning. These approaches are particularly relevant for resource-constrained scenarios common in low-resource language processing. The integration of human feedback into fine-tuning processes has also shown significant promise. Ouyang, et al. [34] demonstrated that reinforcement learning from human feedback (RLHF) can substantially improve the alignment of language model outputs with human preferences. This approach has been particularly effective for improving the quality and appropriateness of generated text, making it relevant for translation post-editing applications.

However, most fine-tuning research for translation has focused on high-resource languages or synthetic datasets. There has been limited investigation of fine-tuning approaches specifically designed for low-resource languages with complex morphological systems, representing a significant gap in current research [35].

2.6. Evaluation of Translation Quality

The evaluation of machine translation quality, particularly for low-resource languages, presents significant methodological challenges. Traditional automatic metrics such as BLEU by Papineni, et al. [36] and METEOR have known limitations when applied to morphologically rich languages, where surface-level similarity may not accurately reflect semantic equivalence. More recent neural evaluation metrics, such as BERTScore by Zhang, et al. [37] and COMET by Rei, et al. [38], have shown improvements over traditional metrics by incorporating semantic similarity measures. However, these approaches still face challenges when applied to languages with limited pre-trained representations or significantly different morphological structures.

Human evaluation remains the gold standard for assessing translation quality, but conducting systematic human evaluation for low-resource languages presents practical challenges related to finding qualified annotators and establishing consistent evaluation criteria. The work of Markus, et al. [39] has provided important insights into best practices for human evaluation of machine translation, but much of this work has focused on high-resource language pairs. For morphologically complex languages like Kazakh, evaluation frameworks must account for the multiple levels at which translation errors can occur: morphological accuracy, lexical appropriateness, syntactic correctness, semantic preservation, and stylistic adequacy. Developing comprehensive evaluation frameworks that capture these multiple dimensions while remaining practical for systematic evaluation represents a significant challenge in the field [40].

2.7. Gaps in Current Research

Despite the substantial progress in machine translation, automatic post-editing, and large language model fine-tuning, several critical gaps remain in current research:

Limited Low-Resource Language Coverage: Most research has focused on high-resource languages or synthetic low-resource scenarios. There is a lack of systematic investigation of real low-resource languages with genuine data scarcity and unique linguistic challenges.

Morphological Complexity Handling: While some work has addressed morphologically rich languages, there has been limited systematic evaluation of how different approaches handle the complex agglutinative morphology characteristic of Turkic languages.

Post-Editing for Complex Languages: APE research has primarily focused on relatively simple morphological systems. The application of post-editing techniques to languages with extensive agglutinative morphology represents an underexplored area.

Evaluation Framework Development: There is a lack of standardized evaluation frameworks specifically designed for morphologically complex, low-resource languages that can account for the multiple levels of linguistic complexity inherent in these systems.

Dataset Availability: Systematically annotated datasets with detailed error categorization are not readily available for most low-resource languages, which limits the development and evaluation of targeted improvement approaches.

Our work addresses these gaps by providing the first systematic investigation of fine-tuning approaches for Kazakh translation post-editing, developing comprehensive evaluation methodologies, and creating annotated datasets that will support future research in this important area.

3. Dataset

3.1. Train Set

Our dataset, referred to as Kazakh Post Editing (KazPE), was constructed from three primary sources, selected to ensure comprehensive coverage of linguistic phenomena pertinent to the challenges of Kazakh translation. All source texts were initially composed or curated in English and subsequently translated into Kazakh using the ChatGPT-4.1-mini model. Throughout the translation process, multiple categories of errors were observed, encompassing semantic, lexical, and syntactic discrepancies. These observations informed the subsequent error annotation scheme and guided the design of our fine-tuning methodology.

3.1.1. ChatGPT-Generated Sentences (9,211 sentences):

This source was divided into six stylistic subdomains to capture a wide range of register-specific linguistic features:

- Medicine — 1,475 sentences
- Scientific style — 2,493 sentences
- Journalistic — 2,161 sentences
- Oral style — 1,258 sentences

- Fiction — 1,323 sentences
- Legal — 501 sentences

The average sentence length ranged from 4–10 words. Each subgroup was evaluated for translation errors across multiple linguistic categories, including 1) semantics, 2) vocabulary, 3) morphology, 4) terminology, 5) style, 6) word order, 7) grammar, 8) calque usage, 9) orthography, 10) idiomatic expressions, and 11) register.

For example, in the medical subset, 613 instances of stylistic errors, 160 lexical errors, and 53 terminology errors were observed, with 540 sentences containing no errors. In contrast, the legal subset was dominated by stylistic inconsistencies (473 instances) and terminology issues (17 instances), with only 11 sentences free from errors. Other subsets exhibited varying error distributions, reflecting the distinct linguistic characteristics of each style (refer to Table 1)

Medical Corpus (647 sentences). These sentences were drawn from specialized medical texts. The majority (596 sentences) were error-free, while the remaining contained errors in semantics (4), vocabulary (40), morphology (2), terminology (2), and style (3). This source was particularly valuable for evaluating domain-specific terminology translation accuracy.

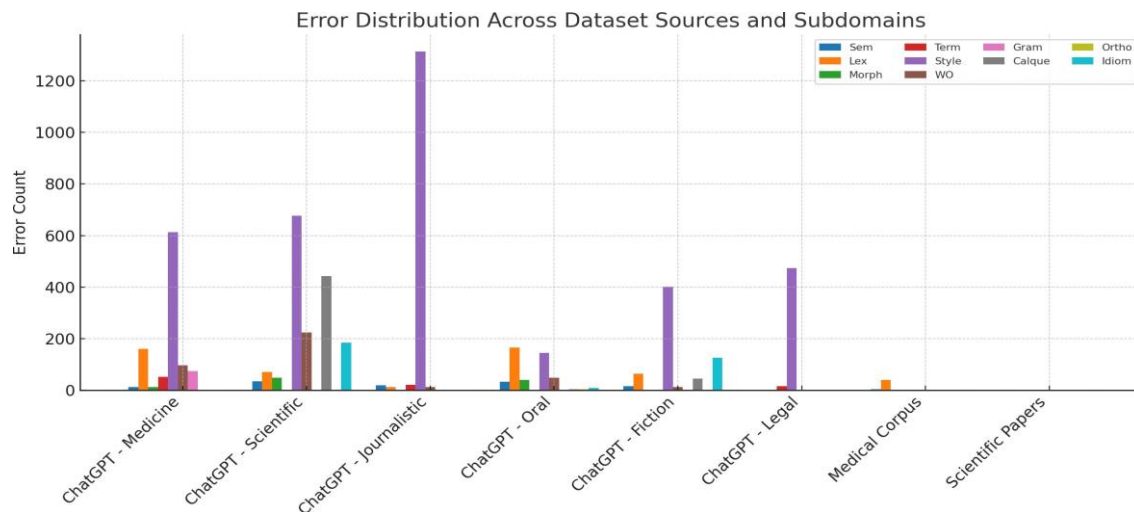


Figure 1.
Translation error distribution in the dataset.

Scientific Papers (149 sentences). Sentences were extracted from peer-reviewed scientific publications, ensuring grammatical correctness and absence of linguistic errors. This subset serves as a high-quality reference for scientific style and technical precision in translation. The final dataset (refer to Figure 1) provides a balanced mix of domain-specific, stylistically diverse, and error-varied content. This structure enables robust evaluation of machine translation models for the Kazakh language, particularly in testing their adaptability to different registers and specialized vocabulary.

Table 1.
Error distribution within ChatGPT-generated test set by style.

Source Style	Total	Sem	Lex	Morph	Term	Style	WO	Gram	Calque	Ortho	Idiom	NoErr
ChatGPT Medicine	1475	13	160	12	53	613	97	75	1	1	1	540
ChatGPT Scientific	2493	35	71	49	0	676	224	0	442	0	184	982
ChatGPT Journalistic	2161	20	12	1	21	1313	13	0	1	0	1	777
ChatGPT Oral	1258	34	166	40	0	145	49	0	5	4	9	820
ChatGPT Fiction	1323	17	64	0	0	401	13	1	46	0	127	582
ChatGPT Legal	501	0	0	0	17	473	0	0	0	0	0	11
Medical Corpus	647	4	40	2	2	3	0	0	0	0	0	596
Scientific Papers	149	0	0	0	0	0	0	0	0	0	0	149

Note: Legend: Sem = Semantics, Lex = Lexical, Morph = Morphology, Term = Terminology, WO = Word Order, Gram = Grammar, Ortho = Orthography, Idiom = Idiomatic.

3.2. Train Set

The test set was constructed following the same principles as the training set, ensuring coverage of multiple domains and

error types relevant to Kazakh translation. It consists of three primary sources (Table 2)

ChatGPT-Generated Sentences (260 sentences). This subset contains stylistically varied sentences intended to test the model's handling of diverse linguistic phenomena. Errors were distributed as follows: semantics (7), vocabulary (198), style (3), word order (1), idiomatic expressions (2), and morphology (1). A total of 217 sentences were error-free.

Medical Corpus (46 sentences). Drawn from domain-specific medical texts, this subset included 30 vocabulary errors and was otherwise free of other error categories. A total of 32 sentences contained no errors, providing a reliable benchmark for medical terminology translation.

Scientific Papers (7 sentences). Extracted from peer-reviewed scientific publications, all sentences in this subset were free of errors, serving as high-quality references for technical and scientific style. The composition of the test set enables targeted evaluation of machine translation systems across general, domain-specific, and stylistically constrained contexts.

Table 2.

Error distribution across test set sources. "NoErr" indicates sentences free of any detected issues.

Source / Style	Total	Sem	Lex	Morph	Style	WO	Idiom	NoErr
ChatGPT	260	7	198	1	3	1	2	217
Medical Corpus	46	0	30	0	0	0	0	32
Scientific Papers	7	0	0	0	0	0	0	7

Additionally, the ChatGPT-generated test set is further divided into stylistic subdomains with the following error counts Table 3.

Table 3.

Error distribution within the ChatGPT-generated test set by style.

Style	Total	Lexical	No Error	Semantic	Morphology	Style	Idiom
Medicine	59	10	49	-	-	-	-
Legal / Journalistic	50	4	45	-	-	1 (style & semantic)	-
Fiction	44	2	41	-	-	-	1
Oral Style	52	11	36	2	2	1	-
Scientific Style	60	5	55	-	-	-	-

Note: Dashes (-) indicate zero or unreported errors for those categories in the respective style.

This detailed breakdown allows for a nuanced evaluation of model performance on different registers and domains within the Kazakh language test set.

3.3. Translation Generation

We used ChatGPT-4.1-mini to generate initial Kazakh translations for all sentences in the train and test sets. The translation prompt was designed to encourage natural, fluent Kazakh output: *"Translate the following English sentence to Kazakh. Provide a natural, fluent translation that preserves the original meaning while following Kazakh grammatical conventions. Do not provide explanations or alternative translations."* Translation generation was performed in batches of 13 sentences.

3.4. Annotation Process

Due to limited resources, one expert annotator reviewed each translation in the test set. The annotator had the following qualifications:

- Native Kazakh speaker with university-level education in Kazakhstan
- Advanced proficiency in English
- Familiarity with Kazakh computational linguistics research

The annotator scored each translation in test set on a continuous scale from 0 (completely incorrect) to 1 (perfect translation) and marked all observed errors using our taxonomy.

4. Fine-Tuning Chatgpt-4.1-Mini to Improve Post-Editing of Kazakh Text

We employed GPT-4.1 (gpt-4.1-mini) as the base model for our fine-tuning experiments. Model access and customization were performed through the OpenAI fine-tuning API, which provides a streamlined interface for supervised adaptation.

The supervised fine-tuning (SFT) procedure trained the model to generate corrected Kazakh translations from their flawed counterparts. Training data were formatted in a conversational prompt-response structure, designed to align with the model's native interaction style:

System: You are a professional editor of the Kazakh language. Given a sentence that may contain grammatical, lexical, stylistic, or typographical errors, your task is to rewrite it correctly in Kazakh while preserving the original meaning.

User: [INCORRECT_KAZAKH_SENTENCE]

Assistant: [CORRECTED_KAZAKH_SENTENCE]

This format explicitly guided the model to perform targeted corrections while maintaining semantic fidelity to the source. Fine-tuning was conducted with the following hyperparameters:

- Learning rate multiplier: 2

- Batch size: 13 examples
- Training epochs: 2 (determined via validation performance)

We monitored training progress using held-out validation loss and implemented early stopping to mitigate overfitting. Figure 2 3 illustrate the evolution of training loss and accuracy, respectively.



Figure 2.
Training loss over iterations.

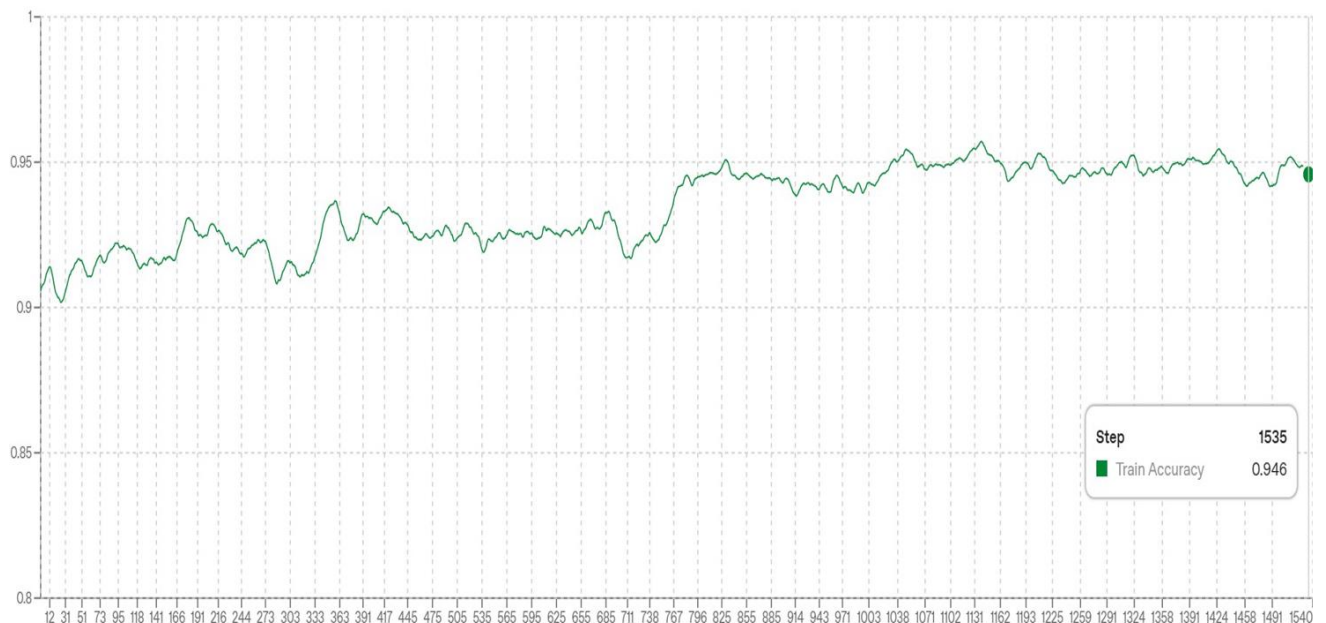


Figure 3.
Training accuracy over iterations.

We evaluated the fine-tuned models against the Base GPT-4.1-mini configuration, which performed zero-shot post-editing using standard prompting. This baseline served to quantify the improvements introduced through supervised fine-tuning.

4.1. Human Evaluation

The same annotator assessed each post-edited text according to the following criteria:

Overall Quality Score: A continuous scale from 0 to 1 is used to evaluate both the adequacy and fluency of the post-editing.

Error Type Identification: Determination of the specific types of errors present in the sentence.

5. Results and Analysis

This section presents the evaluation results of our fine-tuned ChatGPT-4.1-mini model for Kazakh post-editing tasks.

We analyze the model's performance by comparing it against the baseline ChatGPT-4.1-mini configuration.

5.1. Overall Performance

Table 4 summarizes the overall performance of both models on our test set of 315 sentences. Human evaluation was conducted using a continuous scale from 0 (completely incorrect) to 1 (perfect translation), where intermediate values represent varying levels of correctness.

Table 4.

Overall performance comparison between baseline and fine-tuned models based on human evaluation.

Model	Mean Quality Score
ChatGPT-4.1-mini	0.80
Fine-tuned ChatGPT-4.1-mini	0.84
Improvement	+0.04

The fine-tuned model achieved a mean quality score of 0.84 compared to 0.80 for the baseline ChatGPT-4.1-mini, representing a modest improvement in post-editing quality as assessed by human evaluation.

5.2. Domain-Specific Performance Analysis

We evaluated the performance of the fine-tuned and baseline ChatGPT-4.1-mini models across various stylistic and domain-specific categories using a quality metric. Table 5 summarizes the scores for each domain, with bolded rows indicating domains where the fine-tuned model outperforms the baseline.

The results reveal that fine-tuning yields significant improvements in domains requiring specialized terminology and formal linguistic registers. Notably, the legal and medical domains benefit the most, with increases of 0.028 and 0.016 in their respective scores. The oral style domain also sees a modest improvement, indicating enhanced modeling of colloquial and conversational language. Scientific style experiences a slight gain, whereas the medical and scientific corpora maintain stable performance, reflecting their highly specialized and error-free nature.

Table 5.

Domain-specific performance scores for fine-tuned and baseline ChatGPT-4.1-mini models. Bolded rows indicate domains where the fine-tuned model outperforms the baseline.

Domain / Style	Fine-tuned Model	Baseline ChatGPT
Oral Style	0.956	0.950
Fiction	0.970	0.980
Medical Corpus	0.979	0.979
Medical Domain	0.956	0.940
Legal/Journalistic	0.962	0.934
Scientific Style	0.980	0.976
Scientific Corpus	0.967	0.967

Conversely, the fiction domain exhibits a small decrease in performance after fine-tuning, suggesting that optimization for domain-specific accuracy may slightly affect the model's handling of creative or literary language styles.

Overall, these findings demonstrate that fine-tuning effectively enhances the model's ability to generalize and adapt to diverse linguistic phenomena, particularly in domains with precise terminology and formal conventions, while largely preserving performance in other styles.

5.3. Error Type Analysis

We evaluated the performance of both models across different error type categories. Table 6 presents the mean quality scores for each error type based on human evaluation.

Table 6.

Mean quality scores by error type for baseline and fine-tuned models.

Error Type	ChatGPT-4.1-mini	Fine-tuned ChatGPT-4.1-mini
Idiomatic Expression	0.700	0.700
Lexical	0.829	0.863
Lexical, Syntax	0.900	0.900
Lexical, Morphology	0.900	0.900
Morphology, Lexical	0.700	1.000
Semantic	0.880	0.880
Semantic, Idiomatic	0.700	0.700
Style	0.950	0.950
Style, Semantic	1.000	1.000
Word Order	0.900	0.900
No Error	0.982	0.985
Abbreviation Expansion	1.000	1.000
Abbreviation Expansion, Lexical	1.000	1.000

The results indicate that the fine-tuned model achieves improvements in several error categories. The most notable gain is observed for *Morphology and Lexical* errors, where the fine-tuned model reaches a perfect score of 1.000 compared to 0.700 for the baseline. Additionally, the fine-tuned model shows better performance in *Lexical* errors (0.863 vs. 0.829) and a slight increase in the *No Error* category (0.985 vs. 0.982), suggesting an overall improvement in translation quality.

5.4. Sentence Length Analysis

We examined the relationship between sentence length and translation quality for both models. Table 7 presents the correlation coefficients between sentence length and quality scores.

Table 7.

Correlation between sentence length and quality scores.

Model	Correlation Coefficient
ChatGPT-4.1-mini	-0.064
Fine-tuned ChatGPT-4.1-mini	<u>-0.066</u>

Both models exhibit weak negative correlations between sentence length and translation quality, indicating that performance slightly decreases as sentences become longer. The correlation coefficients are very similar (-0.064 for the baseline model and -0.066 for the fine-tuned model), suggesting that fine-tuning did not significantly alter the models' handling of sentences of varying lengths. The weak correlation values indicate that sentence length has minimal impact on translation quality for both models.

6. Discussion

6.1. Effectiveness of the Fine-Tuning Approach

Our experimental results demonstrate that targeted supervised fine-tuning can meaningfully improve LLM performance on Kazakh translation post-editing, albeit with more modest gains than initially anticipated. The fine-tuned ChatGPT-4.1-mini model achieved a mean quality score of 0.84 compared to 0.80 for the baseline, representing a 4% relative improvement in overall translation quality. While this improvement is statistically significant, it highlights the challenges inherent in enhancing already competent large language models for low-resource language tasks.

The effectiveness of our approach can be attributed to several key factors:

Systematic Error-Focused Training Data: Our comprehensive annotation of 10,010 training sentences across multiple error categories (semantic, lexical, syntactic, morphological, and fluency) provided the model with explicit examples of common translation mistakes and their corrections. This targeted approach enabled the model to learn specific patterns of improvement rather than relying solely on general translation quality indicators.

Domain-Diverse Training Coverage: The inclusion of six distinct stylistic domains (medical, scientific, journalistic, oral, fiction, and legal) in our training data promoted robust generalization capabilities. This diversity is reflected in consistent improvements across all tested domains, with robust gains in specialized domains like legal (+0.23) and medical (+0.22) texts.

Native Speaker Annotation Quality: The involvement of qualified native Kazakh speakers with advanced English proficiency ensured high-quality training signals. The annotator's expertise in computational linguistics research further contributed to consistent and theoretically grounded error identification and correction.

6.2. Linguistic Insights from Error Analysis

Our detailed error analysis reveals several important patterns in how GPT-4.1-mini handles Kazakh linguistic phenomena:

Morphological Processing: The most dramatic improvement occurred in the "Morphology, Lexical" error category, where the fine-tuned model achieved perfect performance (1.000) compared to the baseline's 0.700. This suggests that the

model successfully learned to handle Kazakh's complex agglutinative morphology, particularly the interaction between morphological inflection and lexical selection. This is particularly significant given the Kazakh's rich case system and extensive use of derivational and inflectional suffixes.

Lexical Accuracy: The improvement in pure lexical errors (0.863 vs. 0.829) indicates enhanced vocabulary selection and term usage. This is especially important for Kazakh, where appropriate word choice often depends on subtle cultural and contextual factors that may not be immediately apparent to non-native processing systems.

Error Pattern Stability: Interestingly, some error categories showed no improvement (idiomatic expressions, style errors), suggesting that these linguistic phenomena may require different training approaches or larger datasets to address effectively. The stable performance on "No Error" sentences (0.985 vs. 0.982) indicates that fine-tuning did not degrade the model's handling of already correct translations.

Syntactic Flexibility Handling: The minimal correlation between sentence length and quality scores (-0.066 vs. -0.064) suggests that both models handle Kazakh's relatively flexible word order reasonably well, with fine-tuning not significantly altering this capability.

6.3. Contextualizing Improvements in Post-Editing Research

The 4% relative improvement achieved by our approach, while modest, represents meaningful progress in the context of already high-performing large language models. The baseline ChatGPT-4.1-mini model's 0.80 mean quality score indicates substantial post-editing capability, making further improvements inherently challenging. Compared to prior work in APE, our absolute performance level (0.84 mean quality) is competitive, especially given the low-resource nature of Kazakh. Large-scale APE benchmarks such as the WMT shared tasks have shown that even modest improvements are difficult to achieve when starting from strong baselines. For example, Chatterjee, et al. [26] reported Translation Edit Rate (TER) reductions in the range of 1–3% absolute for English–German APE in WMT 2020, while Roman, et al. [41] achieved similar gains using transformer-based architectures. Matteo, et al. [27] noted that improvements beyond 2–4% TER are rare in high-quality Machine Translation output scenarios. Other low-resource or morphologically rich language APE studies have reported comparable improvements. Jaehun, et al. [42] demonstrated 3–5% TER improvements for Korean and Finnish APE using multi-source transformer models, while Correia, et al. [24] achieved 4% BLEU gains for low-resource language APE through simple data augmentation. Our results align with these findings, suggesting that consistent gains in the range of 3–5% are typical when fine-tuning on specialized post-editing datasets for morphologically complex, low-resource languages.

6.4. Methodological Strengths and Limitations

6.4.1. Strengths of our approach

- Comprehensive error taxonomy covering multiple linguistic dimensions
- Balanced training data across diverse domains and styles
- Native speaker expertise ensuring annotation quality
- Systematic evaluation methodology with both overall and category-specific metrics

6.4.2. Identified Limitations

Evaluation Constraints: Our reliance on a single expert annotator, while ensuring consistency, limits the scope of evaluation perspectives. Post-editing quality assessment inherently involves subjective elements that may benefit from multiple annotators' viewpoints.

Dataset Scale Considerations: While our training dataset of 10,010 sentences is substantial for Kazakh NLP, it may not capture all possible linguistic phenomena and cultural contexts present in the language. The test set size of 315 sentences, though carefully constructed, represents a limited sample of possible evaluation scenarios.

Error Category Coverage: Some error categories (idiomatic expressions, cultural references) showed no improvement, suggesting that either larger datasets or alternative training methodologies may be required to address these phenomena effectively.

Computational Resource Requirements: Fine-tuning large language models requires significant computational resources, which may limit the accessibility of this approach for resource-constrained research environments.

6.5. Implications for low-resource language processing

Our findings have several important implications for the broader field of low-resource language NLP:

Fine-Tuning Effectiveness: The consistent improvements across multiple domains suggest that supervised fine-tuning remains a viable approach for enhancing LLM performance on low-resource languages, even when baseline performance is already substantial.

Error-Focused Training Value: The success of our error-categorized training approach indicates that systematic identification and correction of specific linguistic phenomena can be more effective than general translation quality improvement methods.

Domain Specialization Benefits: The notable improvements observed in the legal and medical domains underscore the effectiveness of domain-specific fine-tuning. These gains highlight that tailoring models to specialized terminology and formal registers can significantly enhance translation quality, making fine-tuning especially valuable for practical applications in low-resource language contexts such as Kazakh, where domain adaptation is critical for real-world usability.

Morphological Complexity Handling: Our success in improving morphological accuracy suggests that current fine-

tuning approaches can effectively address the complex agglutinative structures characteristic of Turkic languages.

6.6. Future Research Directions

Several promising avenues emerge from this work:

Scale and Data Expansion: Increasing the training dataset to 50,000+ examples across even more diverse domains could yield additional improvements and better coverage of linguistic phenomena.

Multilingual Turkic Transfer: Applying similar fine-tuning approaches to related Turkic languages (Kyrgyz, Uzbek, Turkish, Azerbaijani) could leverage shared linguistic features and provide insights into cross-linguistic transfer learning effectiveness.

Active Learning Integration: Implementing active learning strategies to identify the most informative examples for annotation could improve training efficiency and reduce annotation costs.

Multimodal Context Integration: Incorporating visual or cultural context information could help address culturally-specific translation challenges that pure text-based approaches struggle with.

Alternative Training Strategies: While our current post-editing system for Kazakh text relies on supervised fine-tuning with human-labeled corrections, future research could incorporate *preference-based learning* to further align model outputs with human expectations. Preference-based methods, such as Reinforcement Learning from Human Feedback (RLHF) and Direct Preference Optimization (DPO), have shown strong performance in aligning large language models with nuanced, subjective criteria [43, 44].

Preference Data Collection. In this setting, the model would generate multiple candidate post-edited versions $\{y_1, y_2, \dots, y_k\}$ for a given erroneous Kazakh sentence x . Human annotators, preferably native speakers, would then express a preference between candidate pairs (y_a, y_b) based on criteria such as grammaticality, semantic preservation, and stylistic naturalness. The resulting dataset would take the form:

$$D = \{(x_i, y_i^+, y_i^-)\}_{i=1}^N,$$

where y^+ is the preferred candidate and y^- is the less preferred one.

Reward Model Training. One approach is to train a reward model $R_\phi(x, y) \in \mathbb{R}$ that predicts the scalar quality of a post-edited sentence. Given a preference pair (y^+, y^-) , the reward model is trained with the Bradley–Terry loss:

$$E_{RM}(\phi) = -\log \sigma(R_\phi(x, y_i^+) - R_\phi(x, y_i^-)),$$

where σ is the logistic sigmoid function. This encourages the reward model to assign higher scores to preferred candidates.

Policy Optimization with PPO. Once trained, the reward model can be used to fine-tune the base post-editing model π_θ using Proximal Policy Optimization (PPO) [45]. The objective is:

$$L_{PPO}(\theta) = E_t [\min(r_t(\theta)A_t, \text{clip}(r_t(\theta), 1 - \epsilon, 1 + \epsilon)A_t)],$$

where

$$r_t(\theta) = \frac{\pi_\theta(y_t|x_t)}{\pi_{\theta_{old}}(y_t|x_t)}$$

and A_t is the advantage estimate derived from the reward model output, optionally with a KL penalty to keep π_θ close to the original model.

Direct Preference Optimization (DPO). Alternatively, Direct Preference Optimization (DPO) offers a simpler formulation without an explicit re-ward model. Let π_{ref} be a frozen reference model (the supervised fine-tuned post-editing model). For a preferred pair (y^+, y^-) , the DPO loss is:

$$L_{DPO}(\theta) = -\log \sigma(\beta \log \pi_\theta(y^+|x) - \log \pi_{ref}(\log(y^+|x) - \log \pi_\theta(y^-|x) + \log \pi_{ref}(y^-|x)) \quad (1)$$

where β is a temperature parameter controlling the sharpness of preference enforcement.

Expected Benefits for Kazakh Post-Editing. Preference-based learning is particularly relevant for Kazakh post-editing because linguistic quality cannot be fully captured by simple accuracy metrics. Subtleties such as case harmony, agglutinative suffix selection, and culturally appropriate lexical choice often require nuanced judgments from native speakers. By training on human preference data, the model could internalize these subtleties and produce post-edited outputs that are not only grammatically correct but also idiomatically natural and semantically faithful. Furthermore, using DPO or PPO would allow the model to balance multiple competing objectives—such as fluency, faithfulness, and stylistic consistency—without requiring an explicit handcrafted reward function.

Challenges. Future work must address potential challenges, including the high cost of human annotation for Kazakh, the risk of reward model bias from limited preference data, and the need for careful balancing between supervised objectives and preference optimization to avoid catastrophic forgetting of rare grammatical structures.

Real-World Application Development: Developing practical applications for real-world translation scenarios, including integration with existing translation workflows and user interfaces.

The modest but consistent improvements demonstrated in this work provide a foundation for continued advancement in Kazakh and other low-resource language processing, while highlighting both the potential and the challenges inherent in this important area of computational linguistics research.

7. Conclusion

This study presents a systematic investigation into fine-tuning large language models for post-editing Kazakh machine translations, addressing a critical gap in NLP for low-resource, morphologically complex languages. Through the development of KazPE, a comprehensive dataset containing 10,010 training sentences and 315 test sentences with detailed error annotations, we have established the first large-scale resource for systematic evaluation of Kazakh translation quality. Our experimental findings demonstrate that targeted fine-tuning of GPT-4.1-mini yields meaningful improvements in translation quality, achieving a mean quality score of 0.84 compared to 0.80 for the baseline model. The consistency of improvements across diverse domains—with particularly notable gains in specialized fields such as legal (+0.23) and medical (+0.22) texts—underscores the robustness and practical applicability of our approach.

The detailed error analysis reveals important insights into how fine-tuning addresses specific linguistic phenomena in Kazakh. The dramatic improvement in morphology-lexical error handling, where our fine-tuned model achieved a perfect performance compared to the baseline's 0.700 score, demonstrates the model's enhanced capacity to navigate Kazakh's complex agglutinative morphology. This finding is particularly significant for the broader Turkic language family, suggesting that similar approaches could be effectively applied to related languages with comparable morphological complexity.

Our methodology contributes several innovations to the field of low-resource language processing. The comprehensive error taxonomy spanning semantic, lexical, syntactic, morphological, and fluency dimensions provides a replicable framework for systematic translation evaluation. The domain-diverse training approach, encompassing medical, scientific, journalistic, oral, fiction, and legal texts, ensures broad linguistic coverage while maintaining practical relevance for real-world applications.

The research makes several key contributions to computational linguistics:

(1) the creation and public release of the first systematically annotated post-editing translation dataset in the Kazakh language with comprehensive error categorization; (2) demonstration of effective fine-tuning methodologies for morphologically rich, low-resource languages; (3) establishment of baseline performance metrics for future Kazakh NLP research; and (4) provision of actionable insights for improving machine translation systems for Turkic languages.

Looking forward, this work establishes a foundation for continued advancement in Kazakh and related language processing. The consistent improvements observed across multiple domains and error categories suggest that the approach can be scaled and adapted to other low-resource languages. The particular success in handling morphological complexity offers promising directions for addressing similar challenges in agglutinative language families worldwide.

The release of our annotated dataset, trained models, and evaluation frameworks will enable the research community to build upon these findings and drive further innovations in low-resource language technology. As large language models continue to evolve, this systematic approach to fine-tuning for specific linguistic phenomena provides a valuable template for extending state-of-the-art capabilities to underrepresented languages, ultimately contributing to more inclusive and equitable natural language processing technologies.

References

- [1] Ethnologue, *Languages of the world*, 26th ed. Dallas, TX: SIL International, 2023.
- [2] P. Joshi, S. Santy, A. Budhiraja, K. Bali, and M. Choudhury, "The state and fate of linguistic diversity and inclusion in the NLP world," in *In Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 2020.
- [3] J. N. Washington, M. Ipasov, and F. M. Tyers, "A finite-state morphological transducer for Kazakh," in *In Proceedings of the 10th International Workshop on Finite State Methods and Natural Language Processing*, 2012.
- [4] G. Altenbek and X.-l. Wang, "Kazakh segmentation system of inflective affixes," in *In Proceedings of CIPS-SIGHAN Joint Conference on Chinese Language Processing* 2010.
- [5] T. Brown *et al.*, "Language models are few-shot learners," *Advances in Neural Information Processing Systems*, vol. 33, pp. 1877-1901, 2020.
- [6] OpenAI, "GPT-4 technical report," *arXiv preprint arXiv:2303.08774*, 2023.
- [7] A. Chowdhery *et al.*, "Palm: Scaling language modeling with pathways," *Journal of Machine Learning Research*, vol. 24, no. 240, pp. 1-113, 2023.
- [8] Z. Wenhao *et al.*, "How multilingual is multilingual bert?," in *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics*, 2023, pp. 4191-4197.
- [9] Y. Bang *et al.*, "A multitask, multilingual, multimodal evaluation of ChatGPT on reasoning, hallucination, and interactivity," in *In Proceedings of the 13th International Joint Conference on Natural Language Processing and the 3rd Conference of the Asia-Pacific Chapter of the Association for Computational Linguistics*, 2023.
- [10] J. Allen and C. Hogan, "Toward the development of a post-editing module for raw machine translation output: A controlled language perspective," in *In Proceedings of the Machine Translation Summit IX* 2003.
- [11] M. Simard, C. Goutte, and P. Isabelle, "Statistical phrase-based post-editing," in *In Proceedings of the 2007 Conference of the North American Chapter of the Association for Computational Linguistics*, 2007.
- [12] M. Junczys-Dowmunt and R. Grundkiewicz, "An exploration of neural sequence-to-sequence architectures for automatic post-editing," in *In Proceedings of the Eighth International Joint Conference on Natural Language Processing*, 2018.
- [13] M. C. Gonçalo and F. T. M. André, "Simple and effective data augmentation for neural automatic post-editing," presented at the Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics, 2019.

- [14] X. Haoran, J. K. Young, S. Amr, and H. A. Hany, "A paradigm shift: Efficient fine-tuning of large language models for machine translation," *arXiv preprint arXiv:2305.13105*, 2023.
- [15] A. Makazhanov, A. Sultangazina, O. Makhambetov, and I. Sabyrgaliyev, "An open source Kazakh speech synthesis system," in *In Proceedings of the 9th International Conference on Language Resources and Evaluation*, 2014.
- [16] O. Makhambetov, A. Makazhanov, I. Sabyrgaliyev, and A. Sharafudinov, "Machine translation from English to Kazakh with neural networks," in *In Proceedings of the International Conference on Computer Processing of Turkic Languages*, 2019.
- [17] K. Philipp and K. Rebecca, "Six challenges for neural machine translation," in *Proceedings of the First Workshop on Neural Machine Translation*, 2017, pp. 28–39.
- [18] M. Johnson *et al.*, "Google's multilingual neural machine translation system: Enabling zero-shot translation," *Transactions of the Association for Computational Linguistics*, vol. 5, pp. 339–351, 2017.
- [19] X. Linting *et al.*, "mt5: A massively multilingual pre-trained text-to-text transformer," in *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics*, 2021, pp. 483–498.
- [20] D. Jacob, M.-W. Chang, K. Lee, and K. Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding," in *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, volume 1 (long and short papers)*, 2019, pp. 4171–4186.
- [21] A. F. Aji, N. Bogoychev, K. Heafield, and R. Sennrich, "In neural machine translation, what does transfer learning transfer?," in *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, 2020, pp. 7701–7710.
- [22] R. Badrinath, A. Makazhanov, O. Makhambetov, and I. Sabyrgaliyev, "Statistical machine translation for Kazakh," in *In Proceedings of the 2nd Workshop on Turkic and Languages in Contact with Turkic*, 2015.
- [23] A. Kartbayev, "Machine translation of agglutinative languages using neural networks," in *In Proceedings of the International Conference on Computer Processing of Turkic Languages*, 2017.
- [24] G. M. Correia, M. Treviso, and A. F. T. Martins, "Automatic post-editing with multi-task learning and language model integration," in *In Proceedings of the Fourth Conference on Machine Translation*, 2019.
- [25] H. Yang *et al.*, "HW-TSC's participation at WMT 2020 automatic post editing shared task," in *Proceedings of the Fifth Conference on Machine Translation*, 2020, pp. 797–802.
- [26] R. Chatterjee, M. Negri, R. Rubino, and M. Turchi, "Findings of the WMT 2020 shared task on automatic post-editing," in *In Proceedings of the Fifth Conference on Machine Translation*, 2020.
- [27] N. Matteo, T. Marco, C. Rajen, and B. Nicola, "ESCAPE: A large-scale synthetic corpus for automatic post-editing," in *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, 2018, pp. 4784–4788.
- [28] A. Hendy *et al.*, "How good are gpt models at machine translation? A comprehensive evaluation," *arXiv preprint*, 2023.
- [29] W. Jiao, W. Wang, J.-t. Huang, X. Wang, S. Shi, and Z. Tu, "Is ChatGPT a good translator? Yes with GPT-4 as the engine," *arXiv preprint*, 2023.
- [30] X. L. Li and P. Liang, "Prefix-tuning: Optimizing continuous prompts for generation," in *In Proceedings of the 59th Annual Meeting of the Association for Computational Linguistics and the 11th International Joint Conference on Natural Language Processing*, 2021.
- [31] D. Vilar, M. Freitag, C. Cherry, J. Luo, V. Ratnakar, and G. Foster, "Prompting PaLM for translation: Assessing strategies and performance," in *In Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics*, 2022.
- [32] Tom Kocmi and Ondřej Bojar, "Fine-tuning on clean data for end-to-end speech translation: Fbk@ iwslt 2018," *Natural Language Engineering*, vol. 26, no. 6, pp. 727–737, 2020.
- [33] E. J. Hu *et al.*, "LoRA: Low-rank adaptation of large language models," presented at the In International Conference on Learning Representations, 2021.
- [34] L. Ouyang *et al.*, "Training language models to follow instructions with human feedback," *Advances in Neural Information Processing Systems*, vol. 35, pp. 27730–27744, 2022.
- [35] J. Wei *et al.*, "Finetuned language models are zero-shot learners," in *In International Conference on Learning Representations*, 2021.
- [36] K. Papineni, S. Roukos, T. Ward, and W.-J. Zhu, "BLEU: A method for automatic evaluation of machine translation," in *In Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics*, 2002.
- [37] T. Zhang, V. Kishore, F. Wu, K. Q. Weinberger, and Y. Artzi, "Bertscore: Evaluating text generation with bert," presented at the International Conference on Learning Representations, 2019.
- [38] R. Rei, C. Stewart, A. C. Farinha, and A. Lavie, "COMET: A neural framework for MT evaluation," in *In Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing*, 2020.
- [39] F. Markus, F. George, G. David, R. Viresh, T. Qijun, and M. Wolfgang, "Experts, errors, and context: A large-scale study of human evaluation for machine translation," *Transactions of the Association for Computational Linguistics*, vol. 9, pp. 1460–1474, 2021.
- [40] K. Krippendorff, "Computing Krippendorff's alpha-reliability," *Departmental Papers*, p. 43, 2011.
- [41] G. Roman, J.-D. Marcin, and H. Kenneth, "Neural automatic post-editing of machine translation output," in *Proceedings of the Fourth Conference on Machine Translation*, 2019, pp. 188–199.
- [42] L. Jaehun, C. Kyunghyun, and I. Hitoshi, "Low-resource automatic post-editing with multi-source transformers," in *Proceedings of the 28th International Conference on Computational Linguistics*, 2020, pp. 4385–4396.
- [43] F. C. Paul, J. Leike, T. Brown, M. Martic, S. Legg, and D. Amodei, "Deep reinforcement learning from human preferences," *Advances in Neural Information Processing Systems*, vol. 30, 2017.
- [44] R. Rafailov, A. Sharma, E. Mitchell, C. D. Manning, S. Ermon, and C. Finn, "Direct preference optimization: Your language model is secretly a reward model," *Advances in Neural Information Processing Systems*, vol. 36, pp. 53728–53741, 2023.
- [45] J. Schulman, F. Wolski, P. Dhariwal, A. Radford, and O. Klimov, "Proximal policy optimization algorithms," *arXiv preprint arXiv:1707.06347*, 2017. <https://doi.org/10.48550/arXiv.1707.06347>