

e-ISSN 3060-5059

DOI: 10.67227



IJTIMOYIY-GUMANITAR SOHADA ILMIY-INNOVATSION TADQIQOTLAR

ILMIY METODIK JURNAL

Scientific and Innovative Research in
Social and Humanitarian Studies

Vol. 3 • Issue 9

2026



Founded in 2024



publishscience.uz



publish_science_journal@mail.ru



PEER REVIEWED • OPEN ACCESS

RUS-INGLIZ YO'NALISHIDA MASHINA TARJIMASI MUAMMOLARI: ZAMONAVIY TARJIMA TIZIMLARINING QIYOSIY TAHLILI

Utayeva Iroda Baxodirovna

Samarqand davlat pedagogika instituti, dekan, professor.

e-mail: utayevairoda@2mail.ru ORCID: <https://orcid.org/0009-0004-7268-8083>

Annotatsiya.

Ushbu maqolada rus-ingliz yo'nalishidagi mashina tarjimasi Yandex, Google, DeepL va ChatGPT tizimlari misolida BLEU, METEOR va POS-BLEU metrikalari asosida tahlil qilinadi. Tadqiqot natijalari leksik, semantik va sintaktik aniqlik darajasida muayyan farqlar mavjudligini ko'rsatib, Yandex tizimi umumiy ko'rsatkichlar bo'yicha eng yuqori natijani namoyon etganini ko'rsatadi. Shuningdek, tadqiqotda avtomatik baholash metrikalarining cheklovlari yoritilib, ayniqsa ilmiy va ixtisoslashtirilgan matnlarni tarjima qilishda ularni ekspert baholashi va post-tahrir bilan birgalikda qo'llash zarurligi ta'kidlanadi.

Kalit so'zlar: mashina tarjimasi, rus-ingliz tarjimasi, tarjima sifati, Yandex, Google, DeepL, ChatGPT, BLEU, METEOR, POS-BLEU, ilmiy tarjima, post-tahrir.

ПРОБЛЕМЫ МАШИННОГО ПЕРЕВОДА В РУССКО-АНГЛИЙСКОМ НАПРАВЛЕНИИ: СРАВНИТЕЛЬНЫЙ АНАЛИЗ СОВРЕМЕННЫХ СИСТЕМ ПЕРЕВОДА

Утаева Ирода Баходировна

Самаркандский государственный педагогический институт, декан, профессор.

Аннотация.

В статье рассматривается машинный перевод с русского языка на английский на примере систем Yandex, Google, DeepL и ChatGPT с использованием метрик BLEU, METEOR и POS-BLEU. Результаты исследования выявляют различия в лексической, семантической и синтаксической точности, при этом Yandex демонстрирует наиболее высокие общие показатели. В исследовании также рассматриваются ограничения автоматических метрик оценки качества перевода и подчёркивается необходимость их сочетания с экспертной оценкой и постредактированием, особенно при переводе научных и специализированных текстов.

Ключевые слова: машинный перевод, русско-английский перевод, качество перевода, Yandex, Google, DeepL, ChatGPT, BLEU, METEOR, POS-BLEU, научный перевод, постредактирование.

PROBLEMS OF MACHINE TRANSLATION IN THE RUSSIAN-ENGLISH DIRECTION: A COMPARATIVE ANALYSIS OF MODERN TRANSLATION SYSTEMS

Utayeva Iroda Baxodirovna

Samarkand State Pedagogical Institute, Dean, Professor.

Abstract.

This article examines Russian-English machine translation using Yandex, Google, DeepL, and ChatGPT and evaluates their performance based on BLEU, METEOR, and POS-BLEU metrics. The findings reveal differences in lexical, semantic, and syntactic accuracy, with Yandex demonstrating the strongest overall performance. The study also highlights the limitations of automatic translation evaluation metrics and emphasizes the importance of combining them with expert assessment and post-editing, particularly when translating scientific and specialized texts.

Keywords: machine translation, Russian-English translation, translation quality, Yandex, Google, DeepL, ChatGPT, BLEU, METEOR, POS-BLEU, scientific translation, post-editing.

Introduction. Machine translation can be conceptualized in at least two interrelated senses. In its narrow sense, the term refers to the computer-mediated process of translating a text from a source language into a target language with complete or near-complete automation. In its broader sense, machine translation constitutes an interdisciplinary field of scientific inquiry situated at the intersection of linguistics, mathematics, computer science, and cybernetics. Its primary objective is to develop computational systems capable of performing automated translation in the narrow sense [1]. From a technological perspective, machine translation is an automated procedure in which computational algorithms are employed to generate a target-language text from a source-language input. The conceptual foundations of machine translation emerged in 1947, while one of the first public demonstrations of a machine translation system was conducted at IBM headquarters in 1954. The principal objective of such systems is to reproduce the semantic

content of the source text while taking into account its linguistic and cultural characteristics. Nevertheless, unlike human translators, computational systems do not necessarily possess a comprehensive understanding of meaning, communicative intention, and contextual relations; instead, their output is generated on the basis of computational models, algorithms, linguistic resources, and learned patterns. Consequently, significant challenges arise in the accurate representation of semantic nuances, intonation, contextual meanings, and structurally complex grammatical constructions.

The present study draws on both empirical data and theoretical research in the field of machine translation. The theoretical framework encompasses research on neural machine translation (NMT), artificial neural networks, Transformer architectures, multilingual language models, and automatic translation-quality assessment methods. Particular attention is paid to widely applied evaluation metrics, including BLEU (Bilingual Evaluation Understudy), METEOR (Metric for Evaluation of Translation with Explicit ORdering), and other quantitative indicators used to measure the correspondence between machine-generated translations and reference translations. In addition, contemporary machine translation systems, including Google Translate, DeepL, and Yandex, were examined comparatively in order to identify their capabilities, limitations, and recurrent translation problems.

Literature review. To provide a more systematic and fine-grained assessment of translation quality, a Python-based computational evaluation tool was developed within the framework of the study. The program performs automated sentence-level analysis using BLEU, METEOR, and POS-BLEU metrics. This methodological approach makes it possible to quantify different dimensions of translation performance, identify potentially problematic segments in machine-generated output, and determine recurring linguistic discrepancies. The program subsequently generates an analytical report containing the evaluation results, which can be used as an empirical basis for the further optimization of machine translation models and the improvement of overall translation quality.

Translation, however, cannot be reduced to the mechanical substitution of individual lexical items or syntactic structures from one language with their equivalents in another. It involves the interpretation and reconstruction of contextual meaning, lexical semantics, discourse relations, and semantic ambiguity. One of the most persistent challenges in machine translation is lexical polysemy, since a single lexical unit may activate different meanings depending on its linguistic and situational context. For example, the Russian word *ключ* may denote either a key used to open a door or a spring as a natural source of water. If a machine translation system fails to perform appropriate word-sense disambiguation on the basis of contextual information, it may select an incorrect target-language equivalent and consequently produce a semantically inaccurate translation.

Comparable difficulties arise in the processing of anaphoric and coreferential relations, particularly those involving pronouns. In the Russian sentence «Иван пришел в дом, и он был пуст» (“Ivan entered/came to the house, and it was empty”), the pronoun *он* must be interpreted as referring to *дом* (“house”) rather than *Иван* (“Ivan”). Correct interpretation therefore requires the system to conduct not only sentence-level grammatical analysis but also context-sensitive anaphora resolution and coreference analysis. Such phenomena demonstrate that high-quality machine translation depends on the system's capacity to model semantic and discourse-level relationships extending beyond isolated lexical units.

To address these linguistic challenges, contemporary machine translation research increasingly incorporates methods for phraseological analysis, contextual modelling, word-sense disambiguation, syntactic parsing, anaphora resolution, and discourse-level interpretation. The integration of these mechanisms is essential for improving the semantic adequacy, contextual accuracy, and linguistic naturalness of machine-generated translations, particularly when translating polysemous lexical units, phraseological expressions, referential structures, and syntactically complex constructions.

Discussion. Problems arising from contextual ambiguity require the application of more sophisticated computational methods. One of the most effective approaches is large-scale data-driven training, which enables a machine translation system to process numerous instances of lexical items occurring in different linguistic contexts and to determine their contextually appropriate meanings. For this purpose, large and diverse text corpora representing different genres, registers, and thematic domains are compiled and analyzed, allowing translation systems to reproduce the semantic content of source-language expressions with greater precision. Transformer-based architectures, particularly models such as BERT and GPT, have considerably enhanced contextual language processing and improved the preservation of semantic accuracy. Unlike earlier approaches that process linguistic units in a more restricted or sequential manner, Transformer models analyze relationships among multiple components of a sentence or text, thereby producing richer contextual representations of lexical meaning. Such models can be pretrained on extensive textual datasets and subsequently fine-tuned on parallel or domain-specific data to improve their performance in translation-related tasks.

Recent empirical advances associated with transfer learning and pretrained language models have demonstrated that language-model pretraining can substantially improve performance across a broad range of natural language processing (NLP) tasks. These include sentence-level tasks such as natural language inference and paraphrase identification, in which models predict semantic relationships between sentences by processing them as integrated linguistic units. They also include token-level and span-oriented tasks, such as named entity recognition and question answering, where computational models are required to identify and retrieve more fine-grained linguistic and semantic information [2].

To adapt machine translation systems to heterogeneous types of discourse, researchers develop specialized domain-specific subcorpora. These datasets enable translation models to account more accurately for specialized

terminology, discourse conventions, and stylistic features characteristic of particular professional or academic domains. Furthermore, substantial syntactic and morphological differences between languages require careful system optimization. Translation models must account for cross-linguistic variation in word order, inflectional patterns, grammatical agreement, and other morphosyntactic features, particularly when processing morphologically rich languages.

Considerable progress in language technology has also been achieved at the lower levels of natural language and speech processing. In written-language processing, these developments include text segmentation, lexical analysis, morphosyntactic analysis, and syntactic parsing. In spoken-language processing, major advances have been made in automatic speech recognition, text-to-speech synthesis, and speaker recognition. As a result, numerous language-processing applications have become integrated into everyday communication and information-processing practices. Applications related to written language include spelling and grammar checking, monolingual and cross-lingual information retrieval systems, and online machine translation services. In the field of spoken-language processing, commonly used technologies include voice-enabled GPS systems, automatic dictation, speech transcription, and the automated indexing of audiovisual materials. This range of applications also demonstrates the increasing integration of spoken and written language modalities, as illustrated by speech-to-text transcription and text-to-speech conversion [2, pp. 150–155].

In addition to lexical and contextual ambiguity, one of the most challenging issues in machine translation is the adequate translation of phraseological units. Phraseological units are relatively fixed lexical combinations whose overall semantic content is determined by the expression as a whole rather than by the literal meanings of its individual constituents. English expressions such as “it’s high time,” “take your time,” and “help yourself” cannot always be translated through direct lexical substitution because equivalent meanings in other languages may be expressed through structurally different constructions. From the perspective of semantic compositionality, phraseological units may be broadly classified as non-figurative and figurative expressions. Non-figurative phraseological combinations largely retain the lexical meanings of their constituent elements, although their internal structure is relatively fixed and does not permit unrestricted substitution. Expressions equivalent to “play a role” and “be of significance,” for example, belong to this category and occur particularly frequently in academic, professional, and administrative discourse.

By contrast, figurative phraseological units are characterized by semantic non-compositionality, meaning that their overall interpretation cannot be derived simply by combining the literal meanings of their individual components. Russian provides examples such as «погнаться за двумя зайцами», «попасть как кур в ошип», and «лиха беда начало», while English contains expressions such as “through thick and thin,” “tooth and nail,” and “hit the nail on the head.” The meanings of these expressions are established through conventional usage and frequently reflect language-specific cultural associations, figurative imagery, and communicative traditions.

The difficulty of translating phraseological units automatically lies primarily in the fact that their individual components undergo a degree of semantic transformation or desemanticization, making literal word-for-word translation inadequate. Their interpretation and translation therefore require consideration of both linguistic context and culture-specific meaning. Depending on the communicative situation and the degree of cross-linguistic correspondence, several translation strategies may be applied, including the use of a phraseological equivalent, phraseological analogue, calque (loan translation), or descriptive translation. The selection of an appropriate strategy depends on the semantic, stylistic, pragmatic, and cultural characteristics of the source expression.

Conventional machine translation systems frequently fail to identify such expressions as integrated semantic units, which can result in literal renderings and consequent semantic distortion. To address this limitation, contemporary neural machine translation (NMT) technologies are trained on large-scale datasets containing numerous examples of idiomatic and phraseological expressions in different contexts. By analyzing patterns of contextual usage and their corresponding translations, neural models can learn to select target-language expressions that are functionally and semantically appropriate rather than relying exclusively on lexical correspondence.

Context-sensitive architectures and large pretrained language models, including BERT- and GPT-based approaches, further contribute to this process by modeling relationships between lexical items and their surrounding linguistic environment. Consequently, an idiomatic expression such as “a piece of cake” can be interpreted according to its figurative meaning and rendered through a semantically equivalent target-language expression, such as the Russian «проще простого», rather than translated literally. Such context-aware and meaning-oriented translation mechanisms enhance semantic equivalence, reduce inappropriate literal translation, and improve the overall adequacy and naturalness of machine-generated translations.

Technical constraints also play a substantial role in determining machine translation quality. Contemporary machine translation encompasses several types of systems, including fully automated machine translation, machine-assisted translation systems (MAT and HAMT), and Translation Memory (TM) systems. Their performance varies according to the degree of human involvement in the translation process, although translation accuracy in each system is influenced by a range of linguistic and technological factors. Translation Memory systems, which are extensively employed by professional translators through software such as SDL Trados Studio, memoQ, Wordfast, OmegaT, and Across, store previously translated sentences, phrases, or text segments and retrieve them for subsequent reuse. This functionality reduces translators’ workload and contributes to greater translation consistency and quality by relying on previously validated translation units.

The transition to Neural Machine Translation (NMT) marked an important stage in the development of translation technologies because it enabled machine translation systems to process broader contextual information.

Early NMT models were primarily based on encoder-decoder architectures incorporating recurrent neural networks. However, the introduction of the Transformer architecture in 2017 significantly enhanced translation performance through the application of the attention mechanism, which enabled models to establish more effective relationships among different elements of the input sequence. Compared with earlier machine translation paradigms, NMT substantially reduced problems associated with textual coherence and contextual continuity. Nevertheless, despite these technological advances, fully automated systems such as Google Translate and comparable services continue to exhibit several limitations. These include difficulties in processing specialized and domain-specific discourse, insufficient parallel corpora for certain language pairs, and domain mismatch. Domain mismatch occurs when the thematic, genre-specific, terminological, or stylistic characteristics of the data used to train a machine translation model differ from those of the source or target texts encountered during practical application.

In developing and improving its machine translation system, Yandex has implemented several methodological and computational approaches designed to enhance translation accuracy, particularly for Russian-language input. One such approach involves the application of morphological transformations. The principal advantage of this technique lies in its capacity to transform Russian sentences into representations that are more suitable for computational translation. Russian is characterized by a highly developed inflectional morphology, which frequently creates difficulties for machine translation because grammatical information is encoded through numerous word forms, endings, and morphological variations. Morphological transformation techniques help normalize or restructure these forms, thereby reducing morphological complexity and improving the accuracy of the resulting translations.

Another approach employed by Yandex is the Operation Sequence Model (OSM). This model integrates reordering operations and lexical translation decisions into sequences composed of minimal translation units. From a computational perspective, this approach facilitates the modeling of syntactic structure and word-order transformations, which is particularly important when translating between languages characterized by substantially different grammatical and syntactic systems. Consequently, the Operation Sequence Model addresses one of the fundamental challenges in machine translation: structural divergence between source and target languages. In addition, Yandex has applied the concept of transfemes, whereby words are transformed into more structured translation units. Such representations can reduce the complexity associated with translation between languages that differ considerably in their morphological, grammatical, and syntactic organization. Experimental evaluations conducted using the WMT13 and WMT14 test sets demonstrated that these methodological improvements contributed to higher translation accuracy, as reflected in increased BLEU scores.

Another significant technological platform in the field of machine translation is DeepL, developed by the German company of the same name, which specializes in artificial intelligence-based machine translation technologies. In a 2024 study, Linglin Li conducted a comparative evaluation of translation quality across DeepL, Google Translate, and Microsoft Translator. The systems were assessed according to several translation-quality parameters, including accuracy, fluency, and naturalness, while semantic similarity measurement was employed as one of the principal evaluation strategies.

One of the major findings of the study was that DeepL achieved a reported translation accuracy of 94%, substantially exceeding the 86% result obtained by Google Translate. The evaluation incorporated different textual corpora and assessed overall translation accuracy without specifying a particular language pair. With respect to translation fluency, DeepL also produced a higher reported result, including a 146% change in sentence length. Furthermore, among the three translation systems examined, DeepL demonstrated particularly strong performance in cultural adaptation, with a reported score of 92%, and in the appropriate rendering of idiomatic expressions, where it achieved 94% [13].

The DeepL translation process can be described as comprising several analytical and computational stages. At the initial stage, textual material is differentiated according to categories such as specialized terminology, general vocabulary, and structurally simple sentences. Such categorization facilitates more precise lexical selection and processing, thereby contributing to improved translation quality. At the subsequent stage, the system applies part-of-speech tagging and syntactic analysis to support the grammatical accuracy of the target text. Errors occurring at this level may result in inappropriate lexical choices, syntactic distortions, or incorrect interpretations of contextual relationships; therefore, accurate processing of syntactic structures and contextual information constitutes an essential component of high-quality machine translation.

For translation-quality assessment, DeepL-related evaluation procedures may incorporate semantic similarity measures that enable systematic comparison between the source text, machine-generated output, and alternative translations. Such evaluation makes it possible to determine the extent to which semantic content is preserved across source and target texts. The findings of the comparative study indicate that DeepL performs particularly well in terms of translation accuracy and the processing of stylistically complex texts. At the same time, the study reported that Google Translate demonstrated a slight advantage in certain aspects of contextual and cultural adaptation. These findings suggest that, despite substantial advances in neural machine translation, further research is required to improve contextual interpretation, domain adaptation, culturally appropriate rendering, and the overall semantic reliability of machine-generated translations.

Conclusion. Thus, the effective development of machine translation systems requires the formalization of linguistic knowledge and the construction of models capable of representing language at different structural levels. This makes machine translation an inherently interdisciplinary research domain, integrating theoretical and applied linguistics with computational modelling and artificial intelligence. Particular attention must be devoted to the

adaptation and standardization of linguistic terminology, as well as to the formal representation of lexical, morphological, syntactic, and semantic units. These issues become especially significant when processing languages characterized by rich morphological systems, extensive inflection, flexible word order, and complex syntactic structures.

REFERENCES

1. Бакуменко Я. Д., Таджибова А. Н. Современные проблемы и решения в области машинного перевода // *Успехи кибернетики*. – 2025. – Т. 6, № 2. – С. 47–59. – URL: https://ru.jcyb.ru/nisii_tech/article/view/405 (ru.jcyb.ru)
2. *Net.lang: на пути к многоязычному киберпространству* / пер. с англ.; ред. перевода Е. И. Кузьмин, А. В. Паршакова. – Москва: Межрегиональный центр библиотечного сотрудничества, 2014. – 464 с. – URL: <https://www.vestarchive.ru/component/content/article/188-2010-09-12-12-59-28/2861-1netlang-na-pyti-k-mnogoyazychnomy-kiberprostranstvyr-sbornik-analiticheskikh-materialov.html> (vestarchive.ru)
3. Jawahar G., Sagot B., Seddah D. What Does BERT Learn about the Structure of Language? // *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*. – Florence: Association for Computational Linguistics, 2019. – P. 3651–3657. – DOI: 10.18653/v1/P19-1356. – URL: <https://aclanthology.org/P19-1356/> (aclanthology.org)
4. Devlin J., Chang M.-W., Lee K., Toutanova K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding // *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*. – Minneapolis: Association for Computational Linguistics, 2019. – P. 4171–4186. – DOI: 10.18653/v1/N19-1423. – URL: <https://aclanthology.org/N19-1423/> (aclanthology.org)
5. Rogers A., Kovaleva O., Rumshisky A. A Primer in BERTology: What We Know About How BERT Works // *Transactions of the Association for Computational Linguistics*. – 2020. – Vol. 8. – P. 842–866. – DOI: 10.1162/tacl_a_00349. – URL: <https://aclanthology.org/2020.tacl-1.54/> (aclanthology.org)
6. Shukla A., Bansal C., Badhe S., Ranjan M., Chandra R. An Evaluation of Google Translate for Sanskrit to English Translation via Sentiment and Semantic Analysis // *Natural Language Processing Journal*. – 2023. – Vol. 4. – Art. 100025. – DOI: 10.1016/j.nlp.2023.100025. – URL: <https://doi.org/10.1016/j.nlp.2023.100025> (ScienceDirect)
7. Caswell I. 110 New Languages Are Coming to Google Translate // *Google Blog*. – 2024. – URL: <https://blog.google/products-and-platforms/products/translate/google-translate-new-languages-2024/> (blog.google)
8. Fan A., Bhosale S., Schwenk H. et al. Beyond English-Centric Multilingual Machine Translation // *Journal of Machine Learning Research*. – 2021. – Vol. 22, No. 107. – P. 1–48. – URL: <https://www.jmlr.org/papers/v22/20-1307.html> (jmlr.org)
9. Koehn P., Knowles R. Six Challenges for Neural Machine Translation // *Proceedings of the First Workshop on Neural Machine Translation*. – Vancouver: Association for Computational Linguistics, 2017. – P. 28–39. – DOI: 10.18653/v1/W17-3204. – URL: <https://aclanthology.org/W17-3204/> (aclanthology.org)
10. Maruf S., Saleh F., Haffari G. A Survey on Document-Level Neural Machine Translation: Methods and Evaluation // *ACM Computing Surveys*. – 2021. – Vol. 54, No. 2. – Art. 45. – 36 p. – DOI: 10.1145/3441691. – URL: <https://doi.org/10.1145/3441691> (Monash University)
11. Stahlberg F. Neural Machine Translation: A Review and Survey // *arXiv*. – 2019. – arXiv:1912.02047. – URL: <https://arxiv.org/abs/1912.02047> (arXiv)
12. Fitria T. N. Performance of Google Translate, Microsoft Translator, and DeepL Translator: Error Analysis of Translation Result // *Al-Lisan: Jurnal Bahasa*. – 2023. – Vol. 8, No. 2. – P. 115–138. – DOI: 10.30603/al.v8i2.3442. – URL: <https://journal.iaingorontalo.ac.id/index.php/al/article/view/3442> (Journal of IAIN Sultan Amai Gorontalo)
13. Varela-Salinas M. J., Burbat R. Google Translate and DeepL: Breaking Taboos in Translator Training. Observational Study and Analysis // *Ibérica*. – 2023. – No. 45. – P. 243–266. – DOI: 10.17398/2340-2784.45.243. – URL: <https://www.revistaiberica.org/index.php/iberica/article/view/680> (revistaiberica.org)