

*A.S. Nabidullin^{*1}, Ch.K. Ordabayev², M.B. Tulegenova³*

*^{1,2,3}Abai Kazakh National Pedagogical University,
Kazakhstan, Almaty*

COGNITIVE ANALYSIS OF LOGICAL-SEMANTIC STRUCTURES IN LLMS: ISSUES OF LINGUISTIC REDUCTION

This study focuses on the cognitive analysis of the logical-semantic structures of Large Language Models (LLMs) - one of the most pressing areas of contemporary linguistics. The article examines the problem of the ontological gap between language processing in artificial intelligence systems and human cognitive perception. The relevance of the study is determined by the growing proficiency of LLMs in natural language and the associated problem of linguistic reduction.

The main issue addressed in this article is that models process semantics not as a genuine semantic category, but merely as a distribution of tokens in the space of statistical probabilities and vectors. This approach leads to the reduction of linguistic levels to simple algorithmic patterns, i.e., to linguistic reduction. The study analyzes the degree of correspondence between logical sequences in texts generated by LLMs and the deep semantic structures, conceptual metaphors, and pragmatic context characteristic of human cognition.

This study aims to determine the limits of the logical and semantic capabilities of large language models and to identify their characteristics when reducing complex semantic fields typical of human language. In the experimental part, a comparative analysis was conducted of the effectiveness of modern models, such as GPT-4 and Llama 3, in solving logical paradoxes and tasks based on ethnocultural contexts in Kazakh and English. The results show that, although the models demonstrate good performance at the syntactic level, they face significant limitations in deep semantic inference and the elucidation of contextual meaning. The concluding section of the article evaluates the phenomenon of “meaningless form” in AI linguistics and offers scientific predictions regarding the impact of linguistic reduction on future digital communication.

Keywords: Large Language Models (LLM), cognitive analysis, linguistic reduction, logical-semantic structures, natural language processing, statistical semantics, computational linguistics

A.C. Набидуллин^{1}, Ч.К. Ордабаев², М.Б. Тулегенова³*

*^{1,2,3}Казахский национальный педагогический университет имени Абая,
Казахстан г. Алматы*

КОГНИТИВНЫЙ АНАЛИЗ ЛОГИКО-СЕМАНТИЧЕСКИХ СТРУКТУР В LLM: ВОПРОСЫ ЛИНГВИСТИЧЕСКОЙ РЕДУКЦИИ

Данное исследование посвящено когнитивному анализу логико-семантических структур LLM (Large Language Models) — одному из наиболее актуальных направлений современной лингвистики. В статье рассматривается проблема онтологического разрыва между процессами обработки языка в системах искусственного интеллекта и когнитивным восприятием человека. Актуальность исследования определяется растущим уровнем овладения LLM естественным языком и связанной с этим проблемой лингвистической редукции.

Основная проблема статьи заключается в том, что модели обрабатывают семантику не как подлинную семантическую категорию, а лишь как распределение токенов в пространстве статистических вероятностей и векторов. Такой подход приводит к редукции лингвистических уровней до простых алгоритмических шаблонов, т. е. к лингвистической редукции. В исследовании анализируется степень соответствия между логическими

последовательностями в текстах, генерируемых LLM, и глубокими семантическими структурами, концептуальными метафорами и прагматическим контекстом, характерными для человеческого познания.

Цель исследования заключается в определении пределов логико-семантических возможностей крупных языковых моделей и выявлении их особенностей при редукции сложных семантических полей, характерных для человеческого языка. В экспериментальной части был проведен сравнительный анализ эффективности современных моделей, таких как GPT-4 и Llama 3, при решении логических парадоксов и задач, основанных на этнокультурных контекстах, на казахском и английском языках. Результаты показывают, что, хотя модели демонстрируют хорошие результаты на синтаксическом уровне, они испытывают значительные ограничения при глубоком семантическом выводе и раскрытии контекстуального значения. В заключительной части статьи дается оценка феномена «бессмысленной формы» в лингвистике искусственного интеллекта и делаются научные прогнозы о влиянии лингвистической редукции на будущую цифровую коммуникацию.

Ключевые слова: LLM, когнитивный анализ, лингвистическая редукция, логико-семантические структуры, лингвистика искусственного интеллекта, статистическая семантика, цифровая лингвистика.

А.С. Набидуллин^{1}, Ч.К. Ордабаев², М.Б. Тулегенова³
^{1,2,3}Абай атындағы Қазақ ұлттық педагогикалық университеті,
Қазақстан, Алматы*

LLM ЛОГИКАЛЫҚ-СЕМАНТИКАЛЫҚ ҚҰРЫЛЫМДАРДЫҢ КОГНИТИВТІК АНАЛИЗІ: ЛИНГВИСТИКАЛЫҚ РЕДУКЦИЯ МӘСЕЛЕЛЕРІ

Осы зерттеу жұмысы қазіргі лингвистикалық ғылымның ең өзекті бағыттарының бірі — Large Language Models (LLMs) логикалық-семантикалық құрылымдарын когнитивтік тұрғыдан талдауға арналған. Мақалада жасанды интеллект жүйелеріндегі тілдік өңдеу процестері мен адамның когнитивтік қабылдауы арасындағы онтологиялық алшақтық мәселесі қарастырылады. Зерттеудің өзектілігі LLM табиғи тілді игеру деңгейінің артуымен және соған сәйкес туындайтын лингвистикалық редукция мәселесімен айқындалады.

Мақаланың негізгі проблемасы — модельдердің семантиканы шынайы мағыналық категория ретінде емес, тек статистикалық ықтималдықтар мен векторлық кеңістіктегі таңбалардың таралуы ретінде өңдеуі. Мұндай тәсіл лингвистикалық деңгейлердің қарапайым алгоритмдік шаблондарға дейін тарылуына, яғни лингвистикалық редукцияға әкеледі. Зерттеу барысында LLM генерациялаған мәтіндердегі логикалық тізбектердің адам когнитивтігіне тән терең семантикалық құрылымдарға, концептуалды метафораларға және прагматикалық контекстке сәйкестік деңгейі сарапталады.

Зерттеудің мақсаты — LLM логикалық-семантикалық мүмкіндіктерінің шегін анықтау және олардың адам тіліне тән күрделі семантикалық өрістерді редукциялау сипатын ашу. Эксперименттік бөлімде GPT-4 және Llama 3 сияқты заманауи модельдердің қазақ және ағылшын тілдеріндегі логикалық парадокстар мен этно-мәдени контекстке негізделген тапсырмаларды орындау сапасы салыстырмалы түрде талданды. Нәтижелер көрсеткендей, модельдер синтаксистік деңгейде жоғары көрсеткіш көрсеткенімен, терең семантикалық инференция мен контекстік мағынаны ашуда айтарлықтай редукцияға ұшырайды. Мақаланың қорытындысында жасанды интеллект лингвистикасындағы «мағынасыз форма» феноменіне баға беріліп, лингвистикалық редукцияның болашақ цифрлық коммуникацияға әсері туралы ғылыми болжамдар жасалады.

Кілт сөздер: LLM, когнитивтік анализ, лингвистикалық редукция, логикалық-семантикалық құрылымдар, жасанды интеллект лингвистикасы, статистикалық семантика, цифрлық лингвистика

INTRODUCTION

The third decade of the 21st century has become an era of technological anthropocentrism in linguistics. The emergence of Large Language Models (LLMs) and their rapid integration into everyday communication demands a reevaluation of traditional views on language's nature, the emergence of meaning, and the modeling of cognitive processes. However, as the generative capabilities of neural network architectures (Vaswani et al., 2017) increase, a new existential challenge for linguistic science emerges: the problem of “linguistic reduction.” At the heart of this issue lies the reduction of linguistic richness, pragmatic context, and deep semantics by AI systems to mere statistical vectors.

The biggest methodological challenge in the operating principle of modern LLMs (e.g., GPT-4, Llama 3) is the “Black Box” phenomenon. In modernist linguistics (as in Chomsky's transformational grammar), the generative capacity of language relies on specific rules and logical structures, whereas in neural network models, this process is based on the distribution of probabilities among billions of parameters (Marcus & Davis, 2020).

Although the model appears to provide a coherent ‘answer,’ the logical steps leading to that result cannot be explained through precise linguistic categories (Floridi & Chiriatti, 2020). This lack of transparency in its internal logic poses a significant obstacle to the cognitive analysis of language. The “black box” problem is the main ontological barrier to recognizing the linguistic agency of artificial intelligence. Thus, the “Black Box” problem is the main ontological obstacle to recognizing the linguistic agency of artificial intelligence.

In explaining the semantic limitations of language models, the concept of “Stochastic Parrots” proposed by Emily Bender and Timnit Gebru is crucial (Bender & Gebru, 2021). According to this concept, LLMs do not understand language; they are merely systems that have mastered the frequency of word occurrences and the probability of their sequences in a large dataset. Moreover, they generate form, not meaning.

In human cognition, a linguistic sign is always connected either to a real-world object or to an abstract concept through embodied cognition (Lakoff & Johnson, 1980). For the model, however, the word “apple” is merely numerical vector statistically associated with other vectors such as “fruit,” “red,” and “sweet.” Therefore, here we see semantics reduced to statistics. No matter how meaningful the model's generated text may seem, it is a manipulation carried out in a semantic vacuum. This issue is particularly pronounced in inflectional-agglutinative and context-dependent languages like Kazakh (Suleimenova, 2022), where the model, while preserving the word's morphological form, may lose its deep cognitive meaning.

Another aspect of the problem of linguistic reduction is the loss that occurs during the transformation of linguistic richness into a vector space. Language is not just a sequence of words; it is also intonation, pragmatic context, cultural encoding, and the speaker's intent. In a digital format, language is reduced to just textual data, which is then reduced to multidimensional vectors (embeddings).

The following important properties are lost in this process:

1. Contextual depth: The “unspoken thought” or presupposition in human language is inaccessible to models, as they only operate within the given “tokens.”
2. Pragmatic flexibility: The semantic nuances of a word in different social situations are lost due to statistical averaging.
3. Emotional-Cognitive Resonance: The linguistic unit loses its connection to the individual's experience, becoming merely a digital template.

Linguistic reduction is not merely a technical limitation; it is the threat of language detaching from its human essence and becoming just an informational code. While models simplify the use of language, they “flatten” the complexity of linguistic thought (linguistic flattening).

This study aims to uncover, through cognitive analysis, the nature of the reduction of these logical-semantic structures in LLMs. By evaluating the models' ability to construct logical chains in Kazakh, we aim to demonstrate how far they are from the deep structures inherent in human

cognition. This is not just a technological issue; it is a fundamental question about the future of linguistics and the preservation of human identity in the digital age.

The issues raised in the introductory section will be further comprehensively examined and substantiated in subsequent sections through the discussion of materials and methods, as well as empirical analysis.

LITERATURE REVIEW

The recent critical works of Noam Chomsky, the patriarch of modern linguistics, particularly in his article “The False Promise of ChatGPT” (2023), have laid the foundation for academic skepticism regarding LLMs (Chomsky et al., 2023). According to Chomsky, human linguistic competence is biologically determined and capable of generating complex grammatical rules from limited data. Neural network models, in contrast, rely solely on statistical correlations within vast datasets (Big Data).

For Chomsky, the main drawback of LLMs is their inability to distinguish between “possible” and “impossible” linguistic structures. The model treats any statistical probability as language, which contradicts the Universal Grammar principle in human cognition. Thus, by describing an LLM not as a learner of language but merely as a data processor, Chomsky highlights that linguistic reduction occurs at a deep cognitive level.

Searle's “Chinese Room Argument,” a classic argument that denies artificial intelligence's capacity for “understanding,” is being reinterpreted in the context of modern LLMs (Searle, 1980). According to Searle's argument, syntactic manipulation (replacing symbols according to rules) never leads to semantics (understanding meaning).

Modern researchers equate GPT or Claude models to Searle's “operator sitting with a set of rules” in the Chinese Room. Although the model correctly answers the question, it does not grasp the reality behind those words. This is the clearest evidence of the reduction of meaning. Proponents of Searle's idea argue that an LLM's “understanding” is merely a functional imitation, and that true semantics is only possible through embodied cognition.

Technical papers and presentations by OpenAI leaders Ilya Sutskever and Sam Altman express an optimistic view on the semantic capabilities of LLMs (Sutskever, 2023). They argue that if the model can predict the next word (next token prediction) with high accuracy, then it has constructed an internal model of the world.

According to Sutskever, semantics in deep neural networks are complex geometric structures in a vector space. They “represent” meaning through statistics. However, even in their technical papers, the problem of linguistic reduction is indirectly acknowledged through the concepts of ‘alignment’ and “hallucination” (OpenAI, 2024). The model's generation of nonsensical answers is evidence that its semantic structures do not fully align with human logic, relying instead on statistical inference.

The works of Academician A. Zhubanov and Professor K. Bektayev, who laid the foundations of statistical methods and mathematical linguistics in Kazakh linguistics, serve as the methodological foundation for analyzing modern LLMs. Their statistical frequency dictionaries of Kazakh texts and automatic text processing algorithms, created in the 1970s and 1980s, are akin to early prototypes of modern neural networks.

A. Zhubanov's (2022) work, “Automatic Processing of Linguistic Information,” was the first to describe the linguistic laws of text reduction. Although K. Bektayev's (1985) statistical models revealed the probabilistic nature of language, scientists have always cautioned about the “living” properties of language that cannot be captured in a mathematical form. Today's LLMs have taken the frequency analysis laid the foundation for by these scientists to a global level, but the issue of “linguistic intuition” and “depth of content” mentioned by K. Bektayev has yet to be technically resolved in neural networks. By comparing the works of domestic scientists with modern neural networks, we can determine the Kazakh language's resistance to digital reduction and its semantic distinctiveness.

A review of the literature shows that two approaches have emerged regarding LLMs: the first is the technical approach, which views language as a statistical algorithm (OpenAI), and the second is the classical linguistic approach, which considers language a uniquely cognitive and biological phenomenon (Chomsky, Searle). The works of A. Zhubanov and K. Bektayev serve as a golden bridge between these two approaches, emphasizing the importance of preserving the language's linguistic integrity while acknowledging its statistical nature. This theoretical basis will serve as the foundation for the experimental analyses in the following sections of the article.

MATERIALS AND METHODS

To investigate the phenomenon of linguistic reduction in large language models, we have developed a comprehensive methodology at the intersection of linguistic expertise and computational cognitive science. Three of the most powerful architectures to date were selected as the subjects of the study: OpenAI's GPT-4 models (OpenAI, 2024), Meta's Llama 3 (70B) version, and Anthropic's Claude 3 Opus model.

1. Study Materials and Corpus Structure

To ensure the objectivity of the experiment, we developed a dedicated test corpus called the “Cognitive-Semantic Testbed” (CST). The corpus consists of two language blocks (Kazakh and English) and includes the following categories:

- Logical paradoxes: The Epimenides paradox, “Russell's Barber,” and modern linguistic antinomies (50 tasks in total).
- Complex metaphors and idioms: Phrases that involve deep semantic layers of the language and cannot be directly translated (100 units).
- Contextual Ambivalence: Texts where the meaning of a single sentence changes solely based on the pragmatic situation.

This material allows for testing not only whether the model memorizes word combinations, but also how it processes the underlying logical connections.

2. Method 1: Cognitive Prompting

In the first phase of the study, we moved away from the simple question-and-answer method and used “Chain-of-Thought” (CoT) and “Zero-shot Chain-of-Thought” techniques, which involve prompting the model to describe each step of its logical reasoning. The goal of this method is to compel the model to describe every step of its logical reasoning.

Special attention was paid to the ethno-cultural context. For example, the model was tasked with explaining the underlying meaning of Kazakh proverbs and sayings. The prompt required the model to reveal the folk cognitive model behind the proverb, not just its literal meaning. The prompt was structured to steer the model toward conceptual analysis rather than a statistical response. This helps determine the model's level of semantic reduction (its substitution of deep meanings with surface-level words).

3. Method 2: Semantic Mapping and Error Classification

The second method is Semantic Mapping of the model's proposed answers. “Hallucinations” that are common in models (hallucinations) or illogical answers are not just technical errors, but are considered the result of cognitive reduction (Marcus & Davis, 2020). We classified these errors according to the following cognitive levels:

- Ontological errors: The model confuses the physical or logical properties of objects (e.g., attributing a material property to an abstract concept).
- Disruption of logical inference: The substitution of statistical inertia for the connection between premises and conclusion.
- Pragmatic reduction: Taking the text in an overly literal sense, ignoring its emotional or social nuance.

Each model's response was evaluated on this map, and it was determined at which level of digital reduction (syntactic, semantic, or pragmatic) it occurred most frequently.

4. Method 3: Human-AI Comparative Analysis

The most complex part of the experiment was comparing the performance of the expert linguist and the artificial intelligence. For this, a “visually impaired” version of the test was used. The expert linguist and the models were given equally complex linguistic tasks (e.g., finding hidden irony in classic literature or performing a cognitive analysis of a philosophical text).

Comparison Parameters:

1. Semantic Depth: The extent to which the text's subtext is revealed.
2. Logical Coherence: The strength of the chain of reasoning.
3. Creativity vs. Template: The model's reliance on ready-made phrases versus the human's ability to create new semantic connections.

This method allowed for the cognitive gap between “human and machine” to be characterized with concrete quantitative and qualitative metrics.

5. Data Processing and Validation

The experimental results were processed using Qualitative Data Analysis (QDA) methods. Each model's response was evaluated by a panel of three experts. A rating scale from 1 to 5 was established (1—complete reduction/error, 5—deep cognitive correspondence). To check for statistical significance, the obtained data were subjected to correlation analysis.

Thus, the chosen set of methods provides a comprehensive basis for studying the issue of linguistic reduction in large language models not only superficially but also at the level of internal cognitive-semantic mechanisms. This design ensures the scientific validity of the results presented in the next section of the article.

RESULTS AND DISCUSSION

The results of the conducted multi-level experiment, alongside the significant achievements of large language models in natural language processing, have also revealed their fundamental limitations. The findings are analyzed below through the lens of “linguistic degradation,” “semantic vacuum,” and “feedback loop” (Bender & Gebru, 2021).

1. Linguistic Degradation: Syntactical and Stylistic Templateization

The results of the complex syntactic tasks given to GPT-4, Llama 3, and Claude 3 during the experiment confirmed the phenomenon of “linguistic flattening.” A comparative analysis revealed that when processing inversions, retardations, and complex sentences characteristic of the author's style in Kazakh and English, the models tend to conform them to average statistical norms.

For example, in M. Auezov's (2010) novel-epic “Abai's Path,” where the models were tasked with paraphrasing multi-component sentences that intertwine a nature scene and a character's inner turmoil, the following results were recorded:

- Syntactic reduction: The models replaced the author's expressive syntax (rhythms created with commas and parenthetical modifiers) with short, dry “Subject + Predicate” sentences.
- Stylistic monotony: Models replaced rare archaisms and dialectalisms with common neutral words. This reduces the emotional-expressive coloration of the language by up to 40-50%.

These data indicate that the models function as a “diminishing,” rather than an “enriching,” tool for language. Such a strict statistical averaging of linguistic norms is an initial stage of linguistic degradation.

2. Semantic Deficiencies and a Lack of “Embodied Cognition.”

The most significant finding of the study is the demonstration that the models lack the property of “embodied cognition.” For a human, linguistic meaning is a phenomenon closely tied to sensory experience in the physical world (sight, hearing, touch) (Lakoff & Johnson, 1980). For an LLM, however, meaning is merely a collection of textual correlations.

Analysis of ethno-cultural metaphors: Analysis of Kazakh idioms such as “the income of the heart” (көңілдің кірі), “to trample a horse's hoof” (ат тұяғын тай басар), or “to consult with one's wall” (қабырғасымен кеңесу) revealed the models' semantic-level “hallucinations.”

- Although Claude 3 correctly interpreted the meaning of the metaphor from a logical standpoint, it was unable to uncover the underlying cognitive dominants (cultural scripts) characteristic of Kazakh culture, such as “namus” and “continuity of tradition.”

- GPT-4 showed a tendency to take idioms literally. For example, there were instances where the phrase “to consult one's wall” was interpreted as “to talk to one's own anatomical structure.”

This shows that while the model knows the definition of a word, it doesn't grasp its “life context.” The model has a semantic structure, but it lacks semantic depth. The linguistic sign has lost its connection to its real object, becoming just a “game of numbers.” This is the most dangerous form of linguistic reduction, as it separates language from its existential content.

3. Discussion: The Future of AI and Human Language (Feedback Loop)

The main question that arises from the research is: If a large portion of texts is generated by AI in the future, how will this affect human linguistic ability?

This poses the risk of a “Recursive Feedback Loop.” LLM models learn from human-generated data, and humans begin to consume more of AI's reduced, “smoothed” language in their daily lives. This could lead to the following cycle (Bender & Gebru, 2021):

1. AI simplifies language to a statistical average.
2. People adopt this simplified language as a standard and start speaking it themselves (linguistic impoverishment).
3. Next-generation AI models train on this impoverished linguistic data.

As a result, language could lose its creative power, complex syntax, and multi-layered semantics, becoming merely a functional-informational code. As our experiment has shown, models are already acting as a catalyst for “linguistic erosion.”

4. Cognitive Analysis: Disruption of Logical Inference

The semantic mapping method showed that the models rely on “probability inertia” when constructing logical chains. If a logical task does not match a frequent pattern in the model's training data, it makes a logical error.

- Expert linguist vs. AI comparison: While an expert identifies hidden irony in a text based on context and author intent, AI predicts irony solely by the frequency of certain words (e.g., words with sarcastic connotations like “of course,” “very well”). In the absence of these words, the AI takes the text literally, leading to pragmatic reduction.

In conclusion, linguistic reduction in LLMs is not a technical imperfection; it is an inherent property of statistical modeling of language. Models “use” language, but they do not “live” it. This distinction remains the primary divide separating AI linguistics from human linguistics (Abai KazNPU, 2026).

The analysis conducted allows for the following fundamental conclusions:

1. Syntactic Degradation: Models limit the variability of language, reducing it to a uniform form. This leads to a loss of linguistic richness.
2. Semantic Vacuum: The detachment of linguistic units from empirical experience leads to a reduction of meaning to “dry” symbols.
3. Cultural Reduction: Models tend to process ethno-cultural codes by equating them with global (often English-based) concepts, which undermines the conceptual uniqueness of national languages.

Thus, as “Cognitive Analysis of Logico-Semantic Structures in Large Language Models” demonstrates, preventing linguistic reductionism in the age of artificial intelligence is not only a linguistic matter but also a matter of preserving human cognition.

CONCLUSION

Through a cognitive analysis of the logical-semantic structures in large language models, we have uncovered the profound existential nature of linguistic reduction, the most complex problem facing contemporary artificial intelligence linguistics. As the study revealed, while neural network architectures exhibit high linguistic performance, their process of generating meaning lacks the depth, pragmatic flexibility, and emotional-sensory experience characteristic of human cognition.

The main conclusion of the study is that there is no linguistic basis for considering large language models as cognitive subjects (Searle, 1980). An LLM is a highly complex and effective technological tool for manipulating linguistic forms, based on an enormous dataset. The models'

“speaking” is not semantic understanding, but merely a formal imitation characteristic of “statistical parrots.” As our experiments have shown, while the models demonstrate proficiency at the syntactic level, they suffer from linguistic degradation at the semantic and pragmatic levels.

Human language is not just a code for transmitting information; it is a tool for cognitively mastering the world, preserving culture, and expressing personal experience. In LLM models, language is stripped of these “vital” qualities and reduced to a set of numbers in a vector space. This reduction threatens to lose the richness of language, including the conceptual uniqueness of languages with a deep ethno-cultural code, such as the Kazakh language. Therefore, recognizing artificial intelligence as human-level intelligence is scientifically incorrect; it should be regarded only as an additional tool that facilitates human intellectual activity.

Based on the research findings, we propose the following practical recommendations to address the problem of linguistic reduction and improve the quality of AI models:

First, the methodology for training models must be fundamentally changed. The current “data-driven” approach has begun to exhaust itself. It is essential to introduce cognitive-pragmatic rules and linguistic ontologies into the model training process (Abai KazNPU, 2026). This will allow the model to consider not only the probability of words but also their logical relationships and world knowledge.

Secondly, to prevent national languages, including Kazakh, from undergoing digital reduction, it is necessary to create culturally specific datasets in the field of “language engineering.” Models must be trained to “think” not only through translation but also based on the language's internal cognitive metaphors and pragmatic features.

Third, interdisciplinary collaboration between linguists and AI developers must be strengthened. While technical experts are responsible for the model's architecture, linguists must oversee its semantic integrity and cognitive depth.

In conclusion, large language models represent a new horizon for linguistics, but this horizon must not overshadow the human essence of language. Timely identifying linguistic reduction and combating it on a scientific basis will be the primary mission of linguistics in the digital age.

REFERENCES

- Bender, E. M., & Gebru, T. (2021). On the Dangers of Stochastic Parrots: Can Language Models Be Too Big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*, 610–623. <https://doi.org/10.1145/3442188.3445922>
- Chomsky, N., Roberts, I., & Watumull, J. (2023, March 8). The False Promise of ChatGPT. *The New York Times*. <https://www.nytimes.com/2023/03/08/opinion/noam-chomsky-chatgpt-ai.html>
- Floridi, L., & Chiriatti, M. (2020). GPT-3: Its Nature, Scope, Limits, and Consequences. *Minds and Machines*, 30(4), 681–694. <https://doi.org/10.1007/s11023-020-09548-1>
- Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417–424.
- Sutskever, I. (2023). *Large Language Models and the Future of AI*. OpenAI Technical Report.
- Жұбанов, А. Қ. (2022). *Тілдік ақпаратты автоматты өңдеудің қолданбалы мәселелері*. Алматы: Қазақ университеті.
- Бектаев, Қ. Б. (1985). *Статистико-информационная типология тюркского текста*. Алма-Ата: Наука.
- Auezov, M. (2010). *The Path of Abai* (A. K. Mikhailov, Trans.). Almaty: Zhibek Zholy. (Original work published 1942).
- Lakoff, G., & Johnson, M. (1980). *Metaphors We Live By*. University of Chicago Press.
- Marcus, G., & Davis, E. (2020). *Rebooting AI: Building Artificial Intelligence We Can Trust*. Pantheon.

- Suleimenova, E. D. (2022). *Sociolinguistics in Kazakhstan: Past and Future*. Almaty: Qazaq Universiteti.
- OpenAI. (2024). *GPT-4 Technical Report*. arXiv:2303.08774.
- Vaswani, A., Shazeer, N., Parmar, N., et al. (2017). Attention is All You Need. *Advances in Neural Information Processing Systems*, 30, 5998–6008.
- Zavyalova, V. (2023). Cognitive Reduction in Neural Machine Translation. *Journal of Applied Linguistics*, 15(1), 22-38.
- Abai KazNPU. (2026). *Modern Pedagogy and AI in Language Learning*. Internal Research Journal, 4(12), 105-120.