



Texts Generated by Artificial Intelligence: Structure and Semantics

Larysa Kravets^{1,*}, Viktória Stefuca², Natalka Libak³, Eleonora Berta⁴, Adalbert Baran⁵

¹ Department of Philology, Ferenc Rakoczi II Transcarpathian Hungarian University, Berehove, Ukraine

² Department of Ukrainian Language and Culture, University of Nyíregyháza, Nyíregyháza, Hungary

³ Department of Philology, Section of Ukrainian Philology, Ferenc Rakoczi II Transcarpathian Hungarian University, Berehove, Ukraine

⁴ Department of Hungarian Philology, Ukrainian-Hungarian Educational and Scientific Institute, State Higher Educational Institution "Uzhhorod National University", Uzhhorod, Ukraine

⁵ Department of Philology, English Philology Section, Ferenc Rakoczi II Transcarpathian Hungarian University, Berehove, Ukraine;
Department of English Language and Culture, University of Nyíregyháza, Nyíregyháza, Hungary

ARTICLE INFO

Article history:

Received 29 June 2026

Received in revised form 25 July 2026

Accepted 29 July 2026

Available online 1 August 2026

Keywords:

Artificial Intelligence; Large Language Models; Generated Text; Ukrainian-Language AI Texts; Language Norms; Syntactic Complexity; Semantic "Hallucinations"; Neutral Style.

ABSTRACT

The article describes texts generated by artificial intelligence, with an emphasis on Ukrainian-language material, which remains understudied. The aim of the study was to identify specific identifying characteristics of AI texts by comparing their semantic, structural, and stylistic parameters with texts written by humans. The study combined systematic analysis and synthesis, a comparative approach, content analysis, and quantitative linguistic methods (Type-Token Ratio, syntactic complexity analysis by T-unit), as well as semantic modeling, fact verification, and semantic-stylistic analysis. It has been proven that Ukrainian-language AI texts formally meet the basic criteria of textuality (cohesion, coherence, articulation), which are implemented through the probabilistic combination of templates rather than the author's cognitive and communicative activity. Typical markers of machine generation have been identified: template composition (introduction – main part (3–5 subtopics) – conclusion), homogeneous paragraphs of "average" length, predominance of direct word order, presence of passive constructions, excessive frequency of formal connectors, structural and lexical monotony, errors in word usage. Semantic analysis revealed a combination of formal correctness with factual "hallucinations," low information density, a predominance of neutral style, superficial, statistically determined expressiveness, and emotional masking. It was concluded that texts generated by artificial intelligence constitute a separate linguistic phenomenon with its own set of characteristics, which requires a special typology and a flexible, updatable analysis methodology. The risks to linguistic norms, speech culture, and information security are emphasized, as well as the need to develop critical competence in users regarding the perception of AI content.

1. Introduction

The era of large language models (further – LLMs), such as GPT-5, Gemini, and Llama, has radically changed approaches to text creation and processing. Generative artificial intelligence has become a full-fledged agent capable of quickly producing coherent, stylistically adapted, and, at first glance, meaningful texts. It would seem that this should help solve many issues, which is what ultimately happened. However, this technological achievement has also given rise to numerous challenges and

* Corresponding author.

E-mail address: kravets.larysa@ujis.in.ua

presented linguistics, which is entirely focused on texts, with a number of new tasks, the main one being the need for a comprehensive study of texts generated by artificial intelligence.

In the short time that artificial intelligence has been actively functioning, linguistics has already accumulated considerable experience in studying the specifics of generated texts. Most of these studies have been conducted on English-language material, due to both the global dominance of English in the Information Technology (further – IT) industry and the fact that large language models are primarily trained on corpora in which English-language data has the largest share. Accordingly, English-generated texts are the main object of stylometric, linguistic, and semantic research, and most methods, classifiers, and indicators are developed with the peculiarities of English grammar and Latin script in mind.

At the same time, Cyrillic texts generated by Artificial Intelligence (further – AI) remain less studied, which creates a certain imbalance. This applies to Ukrainian and other Slavic languages, which, in addition to their graphic peculiarities, differ in morphological models, have a developed inflectional system, specific syntactic structures, and a different punctuation system, which makes the results of studies based on English language material irrelevant for such texts. There is a particularly noticeable lack of research on Ukrainian-language AI texts in terms of semantic accuracy and depth of content; structural and syntactic features of generation in Ukrainian; issues of stylistic variability, individual style, and markers of artificiality.

The Ukrainian database of linguistic research on texts generated by artificial intelligence is still in its infancy, and many theoretical and practical questions – from the genre competence of models to stylometric recognition – remain open. Systematic study of Ukrainian-language generated texts is important not only from the point of view of improving language models that work with Cyrillic, but also for ensuring the reliability and quality of digital content and predicting possible achievements and challenges associated with the use of generated texts.

Another important issue is the entry of generated texts into cultural circulation and their interaction with readers on the level of emotional and imaginative resonance. One possible way to solve this problem is to compare texts generated by artificial intelligence with texts written by humans in semantic, stylistic, and semiotic terms, which makes it possible to identify limitations and gaps associated with the lack of human experience and cultural memory in their “author”. This, in turn, opens up a new direction of research – understanding how artificial intelligence can represent, reproduce, or imitate cultural meanings, whether the text it generates is capable of functioning as a full-fledged participant in cultural dialogue, and how the increase in generated content affects language and the process of human text creation.

The purpose of this article is to identify specific identifying characteristics of Ukrainian-language texts generated by large language models through comprehensive analysis, systematization, and comparison of the semantic and structural specifics of their key properties with texts written by humans.

2. Literature Review

The study of linguistic differences between texts generated by large language models (LLMs) and texts written by humans requires reference to the theoretical foundations of these technologies. One of the fundamental works that helps to understand the generation mechanism is *Claude Elwood Shannon's* article “*A Mathematical Theory of Communication*” [1]. The entire architecture of modern LLMs, despite their neural complexity, is based on a probabilistic approach to text. *Claude Shannon* viewed communication as a process in which the probability of the next symbol (word, token)

appearing depends on the sequence of previous elements. This is the concept of entropy, which is a measure of the unpredictability of an information source.

The mechanism of artificial intelligence generation consists in calculating the probable next token. Unlike human cognitive processes, generative models do not form opinions or content in a semantic sense, but sequentially predict the most likely next token based on a huge array of training data. This probabilistic modeling is the main mechanism in the creation of any LLM text [1, 2].

Understanding this theoretical basis is important for explaining a key semantic feature of AI-generated texts: the phenomenon of “hallucinations” (the frequent production of syntactically correct but semantically incorrect constructions). Since the model favors the most probable token pattern, it can generate syntactically flawless but semantically incorrect sentences if this probability pattern is stronger than the actual authenticity in the training corpus [3]. *Claude Shannon’s* work serves as a theoretical basis for understanding how generative models create text, which is important for comparing their structural and semantic features with texts that result from cognitive-semantic intent and corresponding human activity.

The rapid development of large language models (LLMs) has led to an increase in linguists’ interest in a new object and the formation of issues related to the study of texts generated by artificial intelligence, particularly in the field of textual linguistics and linguostylistics. The need for a systematic description of the semantic, structural, and stylistic properties of these texts and the identification of their characteristic features has led to a significant number of studies that point to the emergence of a new direction in linguistics.

Among the works devoted to the problem of the genre and discursive affiliation of generated texts, the article by Brommer et al. [4], which highlights the results of research on the genre competence of generative systems, is noteworthy. Using ChatGPT material, the authors showed how the model reproduces the genre and stylistic features of different types of texts, while revealing limitations in conveying individual style, pragmatic nuances, and social characteristics. The analyzed study is closely intertwined with the didactics of writing and the theory of text genres.

The problem of distinguishing between texts written by humans and those generated by artificial intelligence through stylometric analysis is the focus of numerous researchers [5, 6]. In particular, Przystalski et al. [7] used stylometric analysis to work with short texts and text fragments, which are generally difficult to identify stylistic features. One of the main conclusions of their research was that stylistic differences between texts generated by artificial intelligence and those written by humans are noticeable even in very short fragments. The authors confirm that stylometry is an effective way to identify the characteristics of AI texts. Through their analysis, they were able to identify (a) excessive monotony in syntactic constructions, (b) increased predictability in sentence structure, (c) relatively low frequency of infrequently used words, (d) characteristic punctuation patterns, (e) less individual variability and weak expressiveness. In contrast, texts written by humans are more idiosyncratic, stylistically uneven, and less formalized.

A separate group consists of studies devoted to authorial style and its imitation. In his work, James O’Sullivan argues that artistic texts generated by LLM are quite peculiar in terms of stylistic parameters: they are more monotonous, have predictable rhythms, and less pronounced individual deviations characteristic of human creativity [5].

Empirical data on the differences in style and structure between AI texts and human-written texts can be found in an article by Herbold et al. [8]. It states that large language models, namely GPT-3.5 and GPT-4, are already capable of generating higher-quality texts than human students: “AI models generate significantly higher-quality argumentative essays than the users of an essay-writing online forum frequented by German high-school students across all criteria in our scoring rubric. Writing

styles between humans and generative AI models differ significantly: for instance, the GPT models use more nominalizations and have higher sentence complexity (signaling more complex, ‘scientific’ language), whereas the students make more use of modal and epistemic constructions (which tend to convey speaker attitude)” [8]. This study stimulates discussion about the cultural, educational, and communicative consequences of the widespread use of AI generation, raising the questions: Does communication style change under the influence of these texts? Do they affect human text creation? Is it necessary to review teaching methods with the advent of AI?

A comprehensive and analytical work is the article by Terčon and Dobrovoljc [9], which outlines the main features of texts generated by artificial intelligence at the lexical, grammatical, and discursive levels. The authors concluded that artificial intelligence generates texts with a high level of coherence and formal grammatical correctness, but with reduced semantic richness and a tendency toward formulaic language. They also emphasized: “A new text type with its own specific set of linguistic characteristics is thus beginning to emerge, and with the fast growth of AI-generated text on the internet and in other media, it is becoming crucial to systematically study these characteristics” [9].

In the Ukrainian scientific community, research into texts generated by artificial intelligence is still in its infancy, but in recent years there has been a growing number of studies that directly or indirectly analyze their linguistic and stylistic properties. These issues are most systematically studied in the fields of media linguistics, textual linguistics, stylistics, and information influence research. In particular, N. D. Romankiv and V. M. Yavtushenko analyze the compositional and lexical-grammatical features of media texts generated by artificial intelligence, and also point out the communicative challenges associated with the use of AI in the news and information sphere. They emphasize that such texts mostly have a standardized grammatical structure, are characterized by high cohesion and a certain degree of templatization [10]. These observations coincide with the conclusions of most researchers from the international scientific community. The research by D. M. Telpis and N.V. Kutuz stands out from the rest in that it focuses on stylistic and linguistic deviations characteristic of AI texts. The authors argue that lexical, word-formation, morphological, and syntactic deviations may indicate the involvement of artificial intelligence in the creation of texts, particularly in the context of information and psychological operations in media discourse. A number of linguistic indicators have been proposed that make it possible to identify the artificial origin of a text even in the case of careful editing [11]. Ukrainian researchers are also looking into the possibility of using artificial intelligence in education [12] and linguistic research [13]. Thus, Ukrainian scientific studies focus on several issues: analysis of the linguistic and stylistic features of media texts generated by artificial intelligence; description of linguistic deviations in machine authorship; study of the influence of AI texts on the language of contemporary media discourse; application of artificial intelligence in education; use of AI as a tool for linguistic research.

In general, researchers from different countries agree that texts created by artificial intelligence demonstrate high syntactic correctness and compositional orderliness, but have a limited stylistic range, are semantically unsaturated, and often contain unreliable information.

Current linguistic research on texts generated by artificial intelligence focuses on three key issues:

1. Identifying the linguistic features of AI texts.
2. Determining the ability of LLMs to reproduce genre and style specifics.
3. Developing stylometric methods for recognizing machine generation.

Thus, the question of the linguistic status of texts generated by artificial intelligence is currently the subject of active scientific study, and various approaches – from discursive to statistical-stylometric – allow for a comprehensive understanding of this phenomenon.

3. Methodology

The study used a set of linguistic analysis methods, which made it possible to comprehensively characterize the structural, semantic, and stylistic features of texts generated by artificial intelligence. System analysis and synthesis ensured the generalization of the theoretical basis and the formation of a comprehensive conceptual framework for the study. A comparative method was used to establish fundamental differences between human-created texts and texts generated by language models, particularly in terms of composition, coherence, modality, and stylistic organization.

Content analysis and quantitative linguistic methods were used to evaluate formal parameters, which made it possible to measure structural indicators, including Type–Token Ratio, lexical variability coefficients, and methods of syntactic complexity analysis (in particular, T-unit and other indicators). This made it possible to quantitatively characterize the degree of vocabulary diversity, the level of standardization of syntactic constructions, and typical structural patterns inherent in AI texts.

The assessment of content complexity was carried out using the method of semantic modeling and verification, aimed at determining the depth of thematic development and verifying the factual accuracy of information, which is especially important given the possibility of “hallucinations” in generated texts. In addition, semantic and stylistic analysis was used to identify the characteristics of linguistic expression, modality, emotional tone, and stylistic organization of the texts.

The research material consisted of texts generated by ChatGPT-4, ChatGPT-4o, ChatGPT-5, and ChatGPT-5.1 in scientific and popular science styles.

To ensure the reproducibility of the results, the composition and procedure for forming the research corpus are detailed below. In total, 40 texts generated by artificial intelligence (10 for each model: ChatGPT-4, ChatGPT-4o, ChatGPT-5, ChatGPT-5.1) and 30 texts written by humans belonging to the same functional styles were analyzed.

The AI-text corpus contains 20 scientific (annotations, article fragments, literature reviews) and 20 popular science (blog posts, encyclopedia notes, explanatory articles) samples. Human texts were selected from open academic and popular science Ukrainian-language sources (university repositories, industry portals) in the same genre ratio: 15 scientific and 15 popular science.

The total volume of the AI corpus is approximately 48,000 words (average text length 1,200 words), the human corpus is 45,000 words (average length 1,500 words). Unified thematic prompts were used for generation, covering linguistics, digital communication, social consequences of AI, and the history of the Ukrainian language. Examples of prompts:

- “Write a scientific abstract on the impact of digitalization on linguistic personality”;
- “Create a popular science article explaining how large language models generate text.”

All prompts were formulated in Ukrainian without additional restrictions on style or volume.

The selection of texts was carried out by purposeful sampling (topics relevant to the humanitarian sphere are equally represented in both parts of the corpus). Inclusion criteria: text in Ukrainian, no less than 300 words, no obvious signs of poor-quality generation (for AI) or translation from other languages (for human texts). Samples with excessive fragmentation (less than three connected paragraphs) were excluded.

The evaluation protocol involved a combination of automatic tools (calculation of TTR, T-unit length, frequency of passive constructions and connectors) with manual annotation of structural,

semantic and stylistic features by two independent experts (consistency > 85%). This combination allowed for obtaining both quantitative indicators and qualitative linguistic characteristics, which are discussed below.

4. Results

The main features of a text are considered to be cohesion, coherence, linearity, completeness, articulation, and informativeness. These characteristics were identified based on the analysis of traditional printed texts, when a written message was perceived as static, complete, and orderly. The text appeared as a structure with clearly defined boundaries, consistent logic of presentation, and a relatively stable form. With the development of information and communication technologies at the end of the 20th and beginning of the 21st centuries, the role of text has grown significantly and changed at the same time. Digitalization, which has led to a profound transformation of processes and models, a review of strategies and operations using digital tools, has also affected text, which has become a key unit of digital communication. The digital environment has changed the nature and basic characteristics of text, giving rise to a new phenomenon—digital (electronic) text, which exists in the form of digital data and can be processed by appropriate programs. Under the influence of hypertextuality, multimodality, and interactivity, linearity, completeness, coherence, and integrity have changed significantly. Information and communication technologies have transformed text from a monological, complete unit into a dynamic, networked, and open structure, where the key role is played not by sequence, but rather by connectivity via links.

The emergence of large language models has once again brought the issue of textuality to the fore. Text generated artificially has the same basic characteristics, but implements them using fundamentally different mechanisms. As noted above, generated text does not arise from communicative intention, cognitive understanding, or cultural and social experience, but rather as “the statistical distribution of surface elements in the training data” [14]. This means that its properties are only an imitation of textuality, rather than its creation in a deep semantic dimension. Coherence, in particular, is achieved not through logical or causal relationships, but through predictable language patterns; completeness – through structural typification; informativeness – through the combination of statistically relevant fragments without any guarantee of their actual reliability. However, this does not negate the need for a multifaceted study of such texts, since they are involved in the process of verbal communication, are capable of influencing the addressee, and given that the algorithms of generative language models are constantly being improved.

4.1 Analysis of the structure of texts generated in Ukrainian

The structural analysis aimed to reveal how generative models organize information at different levels: from the syntactic structure of a sentence and interphrase connections to the overall composition of the text. It was found that AI texts are generally well structured at both the micro and macro levels, adhere to the grammatical rules of the Ukrainian language, but are not free from errors of various types.

At the syntactic level, generative models show a tendency to use sentences of medium syntactic complexity, unlike humans, who create sentences of varying lengths and internal organization. Sentences generated in Ukrainian are mostly correctly structured, but they are not free of grammatical errors, the number of which decreases with each subsequent version of the model. In the analyzed Ukrainian-language texts generated by artificial intelligence, sentences with a direct

word order and medium length (T-unit was 15 to 20 words) prevailed, but with different syntactic structures (frequent complex sentences of various types and sentences complicated by adverbial phrases and clarifying insertions), with occasional inversions. No extremely short (one-word, non-expanded) or extremely long (periods) sentences were found. In general, the syntax was characterized by uniformity and stereotypical structures (for example: “X is important because ...”, “It is important to note that ...”, “Mozhna stverduvaty, shcho ...”, “Odnym iz kliuchovykh aspektiv ye ...” etc.), which can be considered a marker of generated text. At the same time, there is a revealing combination of simple, simple complex, and complex sentences, which ensures rhythmic fluency and facilitates the perception of the text.

Texts written by humans are characterized by unevenness and variability in syntactic structure, which manifests itself in the use of both short (one-part, simple extended and unextended sentences) and long sentences (branched structures, periods), complex subordinate constructions with different types of subordination, inversion structures, and rhetorical figures, which together create a unique rhythmic and intonational background.

Another characteristic feature of generated texts is syntactic alienation, which arises from the high frequency of passive verb constructions such as: “robota vykonuietsia (kymos)”, “problema doslidzhuietsia (kymos)”, “tekst rozghliadaietsia (kymos)”, “slova pidbyraiutsia,” etc. The authors assume that this is due to the influence of English-language scientific corpora, in which the frequency of the passive form is quite high. It is likely that artificial intelligence reproduces an English syntactic model that is not characteristic of the Ukrainian language, in which the active type of syntactic construction prevails. However, the presence of passive verb constructions in generated texts may also be caused by numerous Ukrainian-language texts that do not comply with current spelling rules and modern grammatical recommendations.

The stereotypical nature of interphrasal connections (formal connectors) is another marker of generated texts. Among the analyzed linguistic means of connection, various types of linguistic units were recorded, most often means of repeated nomination (lexical (thematic) repetition, semantic equivalents of names (anaphoric pronouns and separate adverbs), means of creating a subject-logical grid of expression (aspect-temporal forms of verbs), linguistic means of organizing expression (interjections *po-pershe*, *po-druhe*, *za povidomlenniam*, *otzhe*), grammatical means of expressing interphrasal connections (word order), means of expressing logical relationships between parts of the text (conjunctions, particles, as well as constructions such as: “ “ “*Yak uzhe zaznachalosia vyshche...*”, “*Tse oznachaie...*” Some of them were used frequently in short fragments (for example: “*Same tomu...*”, “*Same cherez tse...*”, “*Same tse...*”), which reduced the quality of the text and emphasized its artificiality. At the same time, the concentration of various types of formal connectors creates the illusion of coherence, even if the semantic connection is weak.

The above observations were verified statistically. The average Type–Token Ratio for AI texts was 0.82 (SD = 0.04), while for human texts it was 0.88 (SD = 0.05); the difference is statistically significant (t-test for independent samples, $p < 0.001$, Cohen's $d = 1.25$). The length of the T-unit in machine texts turned out to be more homogeneous: an average of 17.8 words (SD = 3.2) versus 22.5 words (SD = 8.6) in human samples ($p < 0.01$). The share of passive verb constructions in the AI corpus reached 11.8%, in the human one – 4.9% ($\chi^2 = 23.4$, $p < 0.001$). The excessive density of formal connectors (“in addition,” “at the same time,” “thus,” etc.) averaged 14.2 occurrences per 1,000 words in the generations versus 8.5 in the human texts. The generalization of these metrics confirms that structural monotony, lexical uniformity, and syntactic alienation are not subjective impressions but systematic quantitative differences.

In AI texts, linguistic means of interphrasal connection that characterize texts written by humans are rare or absent altogether, primarily complex metaphorical transitions, metonymies, phraseological units, periphrases, thematic groups of words, words with opposite semantics, adjectives and verbs derived from them, which create the subject-logical grid of the text [15]. Verbs that form a semantic scheme in human-written texts and ensure the development of content thanks to their lexical meanings and aspectual forms are also limited in their use in generated texts.

In texts generated in Ukrainian, various types of grammatical errors were regularly found within sentences and between parts of sentences: “Ukrainian emerged as a separate language in the Middle Ages, but its development depended on various historical events which took place on the territory of Ukraine,” “Language personality on the Internet is used as people express their unique linguistic style and characteristics in online communication,” (about network jargon) “This term is used for linguistic styles that have emerged in online communication, often characterized by features such as abbreviations, emojis, and creative variations in spelling to compensate for the absence of nonverbal cues in text-based communication, according to David Crystal.” Note that the grammatical error is retained even if the deviant phrase is repeated several times.

The macrostructure of the generated text depends on the user’s query. The model can present concise information in the form of logically ordered lists and short explanations. The AI presents the expanded text in a clear but formulaic structure, which includes an introduction containing the thesis, a main part with three to five subtopics (one paragraph per idea), and a conclusion that summarizes the arguments. The macrostructure of texts written by humans, although similar to the above, always reflects individual logic, with the length of text sections and the number of subtopics in the main part being more variable, and indents and associative transitions possible within each section.

AI text paragraphs also have characteristic features that are evident in their volume, internal organization, and connections with other paragraphs. For comparison, recall that the structure of a paragraph in a text created by a human depends on the author’s intention, individual style, and content complexity. Such a paragraph helps to ensure the thematic grouping of material, the logical development of thought, the organization of argumentation, and rhetorical composition. Generated texts are constructed according to a different principle, determined by the probabilistic nature of LLM.

First of all, the uniformity and “average” length of paragraphs in generated texts are striking. They are mostly the same length (3–5 sentences or 40–80 words), unlike paragraphs written by humans, which can vary in size: from very short (1–2 sentences) to emphasize a point, to very long (8–10 sentences) if the author is developing a complex argument or description. Sentences within a generated paragraph are logically and grammatically connected, but the topic can sometimes fall apart (deviations from the main topic are possible). The structure is uniform and is based mainly on linear coherence, which does not involve the development of micro-topics and does not contain argumentative tension. Transitions between paragraphs lack clear logical markers, and there are sometimes abrupt transitions to new micro-themes. There are no associative connections or individual ways of organizing the material. This is a consequence of the fact that language models operate with local coherence and statistical patterns of the corpus (), but do not possess the global communicative intent and compositional intuition characteristic of humans.

An important characteristic of any text is the frequency and diversity of vocabulary. This is an important marker (Type–Token Ratio), which indicates the ratio of unique words to the total number of words. Human-written texts are characterized by lexical heterogeneity and natural semantic variability, so their Type–Token Ratio (further – TTR) is mostly high, unlike AI texts, which are characterized by lexical monotony, word repetition, and the absence or scarcity of non-trivial, rarely

used words. At the same time, in many texts generated by ChatGPT 5.1, the lexical diversity coefficient ranged between 0.78 and 0.86, which is quite high and characteristic of scientific and analytical texts with a large number of terms.

Another important feature of the text is its compliance with current language norms and the level of literacy. Numerous errors of various types in word usage have been identified in texts generated in Ukrainian. For example: *robota* is used – in context, *pratsia* is correct as the materialized result of some work or activity; a literary work, scientific work, or work of art; *vyhliadaty* – in context, *мати вигляд* is correct; *dozvoliaie* – in context, the correct term is *дає змогу* (allows); *vkliuchaty* – in context, the correct terms are *mistyty*, *okhopliuvaty*; *intentsionalnist* – the correct term is *intentsiinist*, *razom z tym* – the correct term is *vodnochas*; *takym chynom* – the correct term is *otzhe*, etc. Occasionally, there are also spelling and grammatical errors (mainly related to word inflection). Illusory textuality dulls the reader's vigilance, and the specifics of text generation cause regular repetitions of incorrect forms, which leads to a kind of legitimization of them. In the long run, this can negatively affect the Ukrainian literary language and linguistic communication, changing the language habits of native speakers, causing the emergence of new artificial trends in word usage, and undermining linguistic norms.

4.2 Semantic analysis of texts generated by artificial intelligence

It has already been proven that the model does not construct text with internal semantic development, as a human does, but only reproduces generalized knowledge based on statistical patterns [3]. Since the model processes large data sets, the content of the generated text should not be underestimated. It may contain new information and previously unknown facts. The authors believe that the issue of the reliability and depth of such content is important, particularly in Ukrainian-language texts. Semantic analysis helps to evaluate text according to these parameters.

Semantic accuracy and depth of content is a measure that assesses the quality of the information content of a text. During the analysis, the accuracy of the facts presented in the texts was verified using various sources. Numerous cases of inaccurate information presented with absolute syntactic persuasiveness were identified [13]. Cases of using terminology from other fields of science and producing term-like combinations were also found. These observations confirm the presence of numerous “hallucinations” in texts generated in Ukrainian.

An important indicator of the depth of content is generalization and non-trivial conclusions. It has been found that the model generally copes well with generalizations of known facts and ideas (high inductive ability), but it does not interpret or evaluate them as a human does, and it cannot always form truly non-trivial interdisciplinary conclusions or original hypotheses that require thinking and creative rethinking, since it lacks intellectual intentionality.

The generated texts were checked for “wateriness” and information density, which are important semantic parameters for SEO. It was found that the model tends to produce a large number of “filler” phrases () that do not have significant informational load but increase the formal volume of the text and maintain its cohesion. A high proportion of such phrases (low information density) is a semantic marker that reduces the quality of the text, especially in an academic environment.

Within the analysis of modal characteristics and tone of AI texts, the authors observed a tendency to use a small number of modal words (*slid rozghlianuty*, *ymovirno*), which express not the author's intention, but a risk minimization strategy inherent in systems trained on large data corpora with strict ethical filtering. This tendency toward weak modality is a manifestation of the model's built-in desire to minimize the likelihood of categorical statements and, accordingly, to avoid false or

unverified conclusions. [3] This distinguishes generated texts from human-written texts, where modality can be stronger and more expressive, especially in genres that involve persuasion, debate, or emotional engagement.

A characteristic feature of AI texts that affects the quality of discourse is a neutral style of presentation. Ethical alignment, necessary to prevent biased or socially risky statements, at the same time leads to a smoothing of the emotional and evaluative aspects of the statement. As a result, the model generates text that lacks a distinct authorial position, individual style, or passion. Such “defensive neutrality” is effective for informational genres, but significantly reduces the expressive potential of texts focused on artistry, rhetoric, or socially critical statements.

It is also important to note the phenomenon of emotional masking—a case where the model attempts to imitate emotionality by using appropriate lexemes (*radisnyi, tryvozhnyi, zvorushlyvyi*), but does not reproduce the deep emotional logic of the text. The emotional tone here is more a consequence of statistical word selection than the result of experience or conceptual motivation characteristic of human texts. This creates a sense of stylistic superficiality: intense emotional vocabulary can be combined with structural monotony, predictable patterns, and a lack of internal conflict or tension. It is the comparison of these levels—lexical and structural—that makes it possible to identify artificial or overly mechanistic emotionality in the generated text.

Scientific texts generated by artificial intelligence, despite their general orientation toward neutrality, are not devoid of expressiveness, but it is formal rather than conceptual in nature. In particular, phrases that emphasize a certain idea (“*it is important to note,*” “*it is appropriate to use,*” “*one of the main conclusions*”) emphasize the significance of research (“*znachnyi vnesok,*” “*vahomyi vklad,*” “*pryntsyypovo novyi pidkhid,*” “*efektyvne rishennia*”), create the impression of argumentative tension, but are not always supported by real content, since they are selected statistically by the model. The use of intensifiers *such as znachnoiu miroiu, istotno vplyvaie, kliuchovym chynnykom ye, krytychno neobkhidno*, etc. is mostly formulaic and often does not correlate with the actual degree of importance of the statement. However, it creates an illusion of authority and expertise and misleads the reader.

Thus, semantic analysis has shown that texts generated in Ukrainian are correct and formally expressive, but often insufficiently informative, contain factual errors, and lack the characteristics of individual authorial presentation.

5. Discussion

Until now, linguistics has correlated text with categories of signs of linguistic and cultural memory, linguistic and cultural experience, the semiotic space of culture in general (semiotic theory of culture), with the concepts of semantic and emotional energy and the energetic resonance of works of art (linguistic synergetics), with an awareness of the need to take into account the commensurability of sociocultural experience, cognitive bases, linguistic and cultural background knowledge of the author and reader (precedence theory, communicative linguistics). The emergence of texts generated by artificial intelligence has opened up a fundamentally new aspect in the understanding of the traditional object of linguistics. Although at first glance it meets all the criteria of textuality, a generated text differs significantly from a text written by a human. Structurally (syntactically) perfect AI text is not the result of individual linguistic consciousness, and therefore does not reflect cultural and personal experience; its semiotic saturation, semantic content, and cultural marking appear not as manifestations of the author’s creative intention, but as statistically determined reconstructions of patterns accumulated in training corpora. AI texts lack the traditional

mechanism of interaction between “author – cultural space – addressee,” which in texts written by humans is decisive for the production of meanings and emotional and value connotations. Artificial intelligence does not possess the background cultural knowledge that humans have, nor does it have the emotional experience or value-semantic attitudes that shape individual style and authorial intent. As already noted, the text it generates, despite its structural and grammatical consistency, may contain “hallucinations” (unreliable facts), be overly “watery,” or demonstrate structural stereotyping. The question of the entry of such text into cultural circulation and its interaction with the reader at the level of emotional and imaginative resonance remains open and requires careful study.

The contribution of this work is a systematic multi-level comparison of structural, semantic and stylistic parameters of Ukrainian-language texts created by humans and several generations of generative models (ChatGPT-4 – ChatGPT-5.1). Unlike previous studies, which were mostly based on English material or were limited to stylometric indicators, we combined quantitative methods (TTR, T-unit, passive frequency) with manual semantic modeling, fact-checking and pragmatic-stylistic analysis, specially adapted to the grammar and punctuation of the Ukrainian language. This made it possible not only to confirm the known markers (template composition, lexical monotony), but also to identify new, language-specific features: calque of passive constructions not characteristic of the Ukrainian language, systematic errors in word usage (for example, “robota” instead of “pratyia” in a scientific context), as well as the phenomenon of emotional masking, when statistically selected emotional vocabulary is not supported by a real communicative intention. Thus, the work offers an analytical framework focused on morphologically rich Slavic languages, and proves that without taking into account language-specific features, universal AI text detectors remain insufficiently sensitive.

G. Fritz, analyzing coherence in relation to artificial intelligence, quoted N. Chomsky, who, when asked, “What do you think of language models, like GPT-3 from Open AI (provided you’ve heard of it)?” replied, “It’s not a language model. It works just as well for impossible languages as for actual languages. It is therefore refuted, if intended as a language model, by normal scientific criteria. Independently of the refutation, the way it works has no relation to language or cognition generally. Perhaps it’s useful for some purpose, but it seems to tell us nothing about language and cognition generally” [14]. The scientist’s words can be interpreted in different ways: discussing the appropriateness of the name “language models” or analyzing the essence of LLMs. One can agree with the expressed opinion, or one can debate it. However, in our opinion, the statement “it seems to tell us nothing about language and cognition generally” [14] is unconditional.

At the same time, it is also indisputable that generated texts, possessing illusory textuality, are already involved in communicative processes, which necessitates their linguistic study. In particular, there is a need for scientific clarification of the concept of “generated text,” determination of the status of this text, understanding and specification of its textual features, and analysis of its structure and semantics. This is necessary not only for theoretical linguistics, but also for practical areas such as education, journalism, and digital security, where the distinction between texts generated by artificial intelligence and those written by humans is of great importance.

One of the prerequisites for effective analysis and identification of generated texts is the development of a system of criteria that would cover the key semantic, structural, stylistic, and pragmatic characteristics of such texts. In particular, it should take into account:

- a special type of cohesion based on superficial statistical connections;
- formal coherence in the absence of deep semantic logic;

- structural orderliness and predictable compositional organization;
- stylistic standardization and reduced individuality;
- potential semantic instability (“hallucinations”) resulting from probabilistic modeling;
- weak cultural and precedent marking;
- minimal level of authorial intentionality.

The generalization and systematization of these features creates the basis for a comprehensive study of texts generated by artificial intelligence, which will make it possible to solve already pressing problems, as well as predict the impact of such texts on language and society, foresee the prospects for their application and the associated risks. Given the intensity of the development of artificial intelligence models, this seems particularly important: as their capabilities grow, so does the range of their textual characteristics. The authors fully share the opinion of researchers who believe that there is an urgent need for a flexible, updatable methodology that takes into account the new forms of linguistic behavior produced by models of different generations [3, 16].

Among the issues that need to be addressed immediately are textual authenticity and communicative responsibility. Texts generated by artificial intelligence can create an illusion of credibility, which entails certain risks for the reader. Accordingly, this fact reinforces the need for critical reading and verification of information provided by artificial intelligence. Users must be able not only to evaluate content for factual accuracy, but also to understand the specifics of how language models work—their probabilistic nature, lack of responsibility, and limits of generative capabilities. Only with a conscious attitude towards the source of the text can the risks of misinformation, manipulation, or unintentional distortion of knowledge be minimized, which are becoming increasingly relevant in the context of the rapid spread of AI communication.

6. Conclusions

The study shows that the emergence of artificial intelligence as an autonomous participant in communication has led to a fundamental revision of established linguistic ideas about text. The traditional generalized theory of text, based on the ideas of structural completeness, logical sequence, semantic integrity, and cultural marking, is insufficient to describe texts generated by large language models. While formally meeting the criteria of textuality, such texts demonstrate a different nature of origin and functioning: their coherence, informativeness, and composition are the result of statistical calculations rather than the author’s cognitive-communicative activity. This leads to the emergence of specific markers—structural templating, superficial coherence, low cultural saturation and insignificant intentionality, frequent semantic “hallucinations,” and stylistic neutrality.

Compositional, structural, and syntactic analysis of Ukrainian-language AI texts has shown that they have a template structure (introduction – main part – conclusion), generated according to standardized compositional patterns, and demonstrate almost perfect grammatical cohesion in significant parts of the text, but contain grammatical errors.

Semantic analysis reveals that the generated texts are shallow in content, lacking in cultural meaning, and limited in their expression of linguistic and cultural memory and linguistic and cultural experience. They also contain numerous factual errors, which are the result of imitation of content rather than cognitive awareness. They are characterized by semantic and stylistic neutrality, as evidenced by the absence of categorical judgments, evaluative and emotionally expressive vocabulary. The author’s position is not expressed. It has been established that AI texts have low

information density. There is a tendency to use verbose fillers, general phrases that provide formal volume but reduce the semantic depth and informative value of the text.

Thus, compositional and syntactic stereotyping, formal coherence, syntactic alienation, lexical monotony, semantic instability, potential unreliability, semantic and stylistic neutrality, excessive “wateriness,” and weak cultural and precedential marking form a specific profile of text generated by artificial intelligence. The combination of the identified features confirms that the generated text is a new linguistic phenomenon that requires separate conceptual understanding and a clear theoretical definition.

The large-scale introduction of such texts into communication practices gives them real social influence. That is why there is an urgent need to develop a flexible, adaptive, and updatable methodology for their analysis, which will take into account the rapid evolution of language models and new forms of their linguistic behavior. An important component of such a methodology should be a system of reliable criteria for identifying and evaluating generated texts, covering their structural, semantic, pragmatic, and stylistic characteristics.

Issues of textual authenticity and communicative responsibility require special attention. The illusion of credibility created by AI texts through grammatical accuracy and formal logic can mislead the reader by concealing unverified facts, semantic shifts, or incorrect generalizations. Therefore, the development of critical competence in users who are able to adequately evaluate AI text, understand its probabilistic nature and the limits of generative mechanisms, is a necessary condition for the safe integration of AI into the public and academic space.

Another important aspect in the study of texts generated by artificial intelligence is related to the impact on language culture and linguistic norms. The widespread use of such texts in the public sphere, education, media, and everyday communication can contribute to the spread of lexical, grammatical, and stylistic deviations that models reproduce due to the characteristics of their training corpora. Incorrect word usage, calquing of syntactic constructions, excessive standardization, stereotypicality, and a neutral style of presentation can gradually form false ideas in users about the normativity of language units and acceptable ways of constructing utterances.

In general, the results of the study show that Ukrainian-language texts generated by artificial intelligence are not only the subject of linguistic analysis, but also an important factor in changing contemporary communication practices. Further research into this phenomenon is not only of theoretical but also of obvious practical importance, as it will make it possible to predict the impact of AI texts on linguistic norms, speech culture, educational processes, media communication, and digital security, as well as to develop recommendations for their responsible and informed use.

At the same time, it is worth emphasizing the distinction between empirical observations and theoretical generalizations. The obtained quantitative data (reduced TTR, T-unit length, passive voice and connector frequency) are direct results of measurements on the studied corpus, while the conclusions about the potential impact on language norms, education, or digital safety are of the nature of a predictive interpretation based on modern sociolinguistic theory.

A number of limitations of this study should also be acknowledged. First, the corpus covers only two genres and a limited number of models of the ChatGPT family; variability caused by other architectures (Gemini, Llama, DeepSeek) and different prompting strategies remains outside the scope of the analysis. Second, the sample is not representative of the entire set of Ukrainian-language texts, so quantitative indicators should be perceived as an illustration of trends, not as normative standards. Third, the study was conducted in a synchronous manner and does not track the dynamics of change. Future work should expand the genre spectrum, involve larger and more balanced corpora, conduct longitudinal monitoring of the evolution of linguistic characteristics of generations,

and, ideally, develop validated metrics for automatic detection of AI texts in Ukrainian. Only under such conditions can more reliable generalizations be made about the socio-communicative consequences of the distribution of generated content.

Author Contributions

Conceptualization, L.K. and V.S.; methodology, L.K. and N.L.; software, E.B.; validation, V.S., E.B. and A.B.; formal analysis, N.L. and A.B.; investigation, E.B. and A.B.; resources, L.K.; data curation, N.L. and E.B.; writing—original draft preparation, L.K. and V.S.; writing—review and editing, N.L. and A.B.; visualization, E.B.; supervision, V.S. and A.B.; project administration, N.L.; funding acquisition, L.K. All authors have read and agreed to the published version of the manuscript.

Funding

This research received no external funding.

Data Availability Statement

The data supporting the findings of this study are available from the corresponding author upon reasonable request.

Conflicts of Interest

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

Acknowledgement

This research was not funded by any grant.

References

- [1] Shannon, C. E. (1948). A mathematical theory of communication. *Bell System Technical Journal*, 27(3), 379–423. Retrieved from <https://people.math.harvard.edu/~ctm/home/text/others/shannon/entropy/entropy.pdf>
- [2] Brown, T. B., Mann, B., Ryder, N., Subbiah, M., Kaplan, J., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., Askell, A., Agarwal, S., Herbert-Voss, A., Krueger, G., Henighan, T., Child, R., Ramesh, A., Ziegler, D. M., Wu, J., Winter, C., Hesse, C., Chen, M., Sigler, E., Litwin, M., Gray, S., Chess, B., Clark, J., Berner, C., Candlish, S., Radford, A., Sutskever, I., & Amodei, D. (2020). *Language models are few-shot learners*. arXiv. <https://doi.org/10.48550/arXiv.2005.14165>
- [3] Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*. (pp. 610–623). Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>
- [4] Brommer, S., Frick, K., Bursch, A., Rodrigues Crespo, M., & Schwerdtfeger, L. K. (2024). ChatGPT and its text genre competence: An exploratory study. *Weizenbaum Journal of the Digital Society*, 4(4). <https://doi.org/10.34669/wi.wjds/4.4.6>
- [5] O’Sullivan, J. Stylometric comparisons of human versus AI-generated creative writing. *Humanit Soc Sci Commun* 12, 1708 (2025). <https://doi.org/10.1057/s41599-025-05986-3>
- [6] Jaashan, H. M. S., & Bin-Hady, W. R. A. (2025). Stylometric analysis of AI-generated texts: a comparative study of ChatGPT and DeepSeek. *Cogent Arts & Humanities*, 12(1). <https://doi.org/10.1080/23311983.2025.2553162>
- [7] Przystalwski, K., Argasiński, J. K., Grabska-Gradzińska, I., & Ocha, J. K. (2024). Stylometry recognizes human and LLM-generated texts in short samples. *SSRN*. <https://doi.org/10.2139/ssrn.4950812>
- [8] Herbold, S., Hautli-Janisz, A., Heuer, U., Kikteva, Z., & Trautsc, A. (2023). A large-scale comparison of human-written versus ChatGPT-generated essays. *Scientific Reports*, 13, 18617. Retrieved from <https://www.nature.com/articles/s41598-023-45644-9>

- [9] Terčon, L., & Dobrovoljc, K. (2025). *Linguistic characteristics of AI-generated text: A survey*. arXiv. <https://doi.org/10.48550/arXiv.2510.05136>
- [10] Romankiv, N. D., & Yavtushenko, V. M. (2025). Media texts generated by artificial intelligence: Linguistic features and challenges. *Grail of Science*. Retrieved from https://scholar.google.com/citations?view_op=view_citation&hl=en&user=5u5BsRAAAAAAJ&sortby=pubdate&citation_for_view=5u5BsRAAAAAAJ:iH-uZ7U-co4C
- [11] Telpis, D. M., & Kutuz, N. V. (2025). Language deviations as identification of the role of artificial intelligence in the formation of IPsO. Retrieved from https://www.researchgate.net/publication/392507429_Movni_deviacii_ak_identifikacia_rolu_stucnogo_intelektu_u_formuvanni_IPsO
- [12] Melnik, A. V. (2023). The use of artificial intelligence in the educational environment: Potential and challenges. In *Development of pedagogical skills of future teachers in the context of educational transformations: Materials of the III All-Ukrainian Scientific and Practical Conference (April 7, 2023)*. Retrieved from <https://eprints.zu.edu.ua/37171/>
- [13] Kravets, L. V. (2025). Artificial intelligence in linguistic research. In O. O. Malenko (Ed.), *Proceedings of the V International Slavic Studies Conference dedicated to the memory of Saints Cyril and Methodius: "Slavic Studies in the modern world: Methodological and value orientations"* (pp. 52–61). Kharkiv–Shumen: KhNPU; KIFT. Retrieved from <https://dspace.kmf.uz.ua/jspui/handle/123456789/5326>
- [14] Fritz, G. (2022). *Coherence in discourse: A study in dynamic text theory*. Giessen University Library Publications. Retrieved from <https://doi.org/10.22029/ilupub-791>
- [15] Kravets, L. V. (2022). Coherence as a basic textual category. In O. M. Semenog (Ed.), *Text in contemporary research paradigms: Theory and practice* (pp. 48–60). Sumy State Pedagogical University named after A. S. Makarenko. Retrieved from <http://repository.sspu.edu.ua/handle/123456789/12575>
- [16] Floridi, L., & Cowls, J. (2019). A Unified Framework of Five Principles for AI in Society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>