

Explainability and Reliability in Large Language Model Systems: A Survey of SHAP, LIME and Attention Frameworks

Mary Shakkina C; Miranda Lakshmi Travis
PG & Research Department of Computer Science and Artificial Intelligence,
St. Joseph's College of Arts and Science (Autonomous),
Affiliated to Annamalai University, Cuddalore 607001, Tamilnadu, India

Martin Aruldoss
Department of Computer Science, Central University of Tamil Nadu,
Thiruvarur 610005, India

Abstract

Increasing adoption of language models in sensitive sectors such as healthcare, finance, cybersecurity, education, etc., highlights the significance of developing transparent decision support systems. This paper reviews the use of explanation techniques for language and machine learning models, including SHAP, LIME, attention-based explanations and frameworks that transform outputs from models into natural language explanations. It turns out that the existing literature demonstrates that although classical language models demonstrate good predictive and semantic performance, they lack transparency. In this regard, SHAP and LIME help in the detection of feature importance and combination of frameworks lead to explanations being more understandable for users. The literature also outlines some limitations of the existing techniques such as explanation faithfulness, hallucinations, biases, privacy issues, computational complexity and lack of benchmarks for evaluation. The use of explainability in healthcare, finance, cybersecurity, industrial systems and IoT is discussed based on the evidence found in the literature.

Keywords:

Explainable AI; SHAP; LIME; Model Transparency; Trustworthy Machine Learning.

I. Introduction

The rapid progress in artificial intelligence and large language models in the area of generative artificial intelligence has made a significant impact on various sectors including education, healthcare, cybersecurity and industry. This invention was created on the basis of transformer-based architecture and outperformed previous innovations with its great performance in understanding natural languages, creation of new content and helping with decision making (Lee, 2025). Nonetheless, due to the widespread use of this innovation, there arise certain issues related to ethics, biases, reliability and human dependency. More precisely, within the realm of education, specifically foreign language learning, these innovative devices can have both positive and negative effects on the critical thinking skills of students. The systematic literature review revealed that 66.67% of sources pointed to the positive effects of this innovation on developing the skills of students, while another 33.33% of sources indicated the negative consequences of such dependency on technology and underestimating of analytical work of learners (Liu et al., 2025).

Moreover, it has been demonstrated that the assessment of students' progress via such systems is reliable as long as human prompting procedures are used (Jauhiainen & Garagorry Guerra, 2025). Besides educational processes, generative models have begun to be utilized in other fields, among which one may highlight medicine and cybersecurity. As for medicine, it helps to make clinical decisions, discover drugs and medications, showing good potential to improve patients' results while preventing any negative consequences (Ong et al., 2025; Chakraborty et al., 2025). However,

comparative analysis shows that in certain diagnostic situations, expert systems are able to perform better than generative ones, proving their strengths in combination with deterministic systems (Feldman et al., 2025). As for the sphere of emergency care, these systems are able to compete with human doctors or even outdo them in terms of triage accuracy (Arslan et al., 2025).

The use of generative techniques in the sphere of cybersecurity and infrastructure security is another area that is gaining momentum, as researchers focus on their ability to be used in threat detection and in defensive tactics. Using them within critical national infrastructure can bring about greater security and face difficulties linked to trust, privacy, and vulnerability (Yigit et al., 2025; Ferrag et al., 2025). On the contrary, using such models on edge computing systems and on distributed computing systems demonstrates their ability to work efficiently due to localized processing (Nezami et al., 2025). Regarding industry and organizations, generative models affect the practices of workflows, organizational structures, and decision-making strategies. The studies of text mining carried out in various industries indicate the tendency to create policies and guidelines that encourage innovation and ethical responsibility (Jiao et al., 2025).

Even though these models aid in analyzing strategic possibilities, the models demonstrate inconsistencies and biases when analyzed independently or even in a single run (Doshi et al., 2025). Moreover, recent researches have emphasized the importance of fairness, transparency and consistency in generative models and have suggested solutions to tackle the problem of bias and inconsistencies in information management systems (Wei et al., 2025). Another challenge that needs to be taken into consideration is the societal and perceptual impact of these models. Distrust of generative tools could lead to damage to an individual's reputation and evaluation and eventually jeopardize the individual's career opportunities (Kadoma et al., 2025). Despite being revolutionary, these models face certain limitations. The limitations consist of problems such as hallucination, interpretation, benchmarking, and inability to quantify reasoning ability of the model (Roukzroh et al., 2025; McIntosh et al., 2025). Traditional ways of evaluation ignore the dynamic and contextual nature of these models hence the need for an ethical and standardized way of evaluation.

In this paper, the author analyses the notion of generative AI and large language models. The main emphasis in the paper is on the analysis of their applications, pros and cons, constraints and impact on society. This paper is a summary of contemporary literature on the topic published in 2025. The objectives of the Study include but not limited to:

- To analyse the development and potential of generative models and large language models in various industries.
- To analyse their applications in education, health care, cyber security and industries.
- To assess the performance and pros of the mentioned systems in practical applications.
- To analyse the key challenges, such as bias, reliability, ethics issues and perceptual impact of those systems.
- To analyze the limitations of existing approaches for evaluation and benchmarking of the systems.
- To identify the research directions in the field.

This introduction gives the basis to understand the complex issue of generative systems' applications and to realize the importance of overcoming the challenges in order to use them responsibly and effectively. The distribution of publications on the topic is provided in Table.1.

Table.1. Publication Distribution by Research Category

Research Category	No. of Papers	Percentage (%)
LLM Applications	18	28.1
Explainable AI (SHAP/LIME)	15	23.4
Attention Mechanisms	10	15.6
BERT Enhancement Studies	12	18.8
Hybrid Architectures	6	9.4

Based on Table.1, the leading share of publications is related to LLM-based applications at 28.1%. It demonstrates the rising trend of application of contextual learning and generation capabilities in various spheres. Moreover, the explainability-related publications comprise 23.4% due to the rising

concerns regarding transparency and accountability in intelligent systems. The methods of BERT-based and attention-based approaches have taken the leading positions in terms of popularity among the studied articles, as they account for more than a third of the total number of papers. Meanwhile, the architectures based on hybrids and security-related frameworks have relatively low popularity, which provides great possibilities for research in this area.

Table.2. Domain-wise Distribution of Reviewed Studies

Domain	Frequency	Percentage (%)
Healthcare	15	23.4
Education	10	15.6
Cybersecurity	14	21.9
NLP/Text Analytics	12	18.8
Finance	4	6.3
IoT Systems	6	9.4
Industrial Applications	3	4.6

It is evident from Table 2 that the sector of healthcare takes the first place accounting for 23.4% of research works analyzed. Application of intelligent technologies for disease diagnosis, decision support, and patient monitoring in the field of healthcare has brought such result. The next place (21.9%) takes cybersecurity since ensuring protection from different threats becomes more urgent each day. Natural language processing is also highly studied by researchers since correct context analysis is important in numerous cases. On the contrary, finance, industry, and IoT receive relatively little attention, which means there is much room for innovations. Thus, one may state that the major motivation is the necessity to achieve high accuracy and reliability in specific spheres.

II. Review on Explainable AI (XAI) impact on LLMs

With the increasing intricacies associated with today's computational models, more emphasis has been put on transparency, interpretability and trust from users' perspective. Language-based frameworks, though highly advanced in their abilities to generate and process natural text, may be viewed by many as opaque because of their inner workings being hard to perceive. In turn, this problem has generated the need for explainability as one of the main concepts in applications when accountability and responsible decision-making are essential. The current state of research notes that frameworks that support explainability should be developed as they provide opportunities for bridging the gap between the model output and its comprehensibility for people (Bilal et al., 2025; Mumuni & Mumuni, 2025). Another trend that is important for the research domain is connected to the potential of language-based frameworks to solve the problem of interpretability. This means that these frameworks can provide explanations for the meaning of the output produced in terms of narrative presentation of information (López-Pernas et al., 2025). Such an approach is especially valuable for the areas of learning analytics and data science. The importance of explanation of decisions made using machine learning models is especially high in case of such domains as healthcare, finance, and automation. It means that accuracy is not sufficient; however, decisions made by algorithms should be explained and justified. Recent research shows that such interpretability approaches as visualization, attribution, and counterfactuals could improve explainability of machine learning models and, therefore, their trustworthiness and legal compliance (Chittimalla & Potluri, 2025; Kumar et al., 2025). In addition, the combination of predictive and explanatory frameworks leads to increased performance and explainability. For example, in mental health prediction tasks, researchers have demonstrated an increase in the accuracy of predictions combined with a production of meaningful explanations for the outcomes (Ahmed et al., 2025).

Moreover, cybersecurity represents yet another field of research in which decision-making transparency plays an important role. The works conducted recently indicate that when language-based explanations are used together with predictive models, the latter can achieve high accuracy when classifying data with providing insights into the very process of decision-making (Lim et al., 2025). Additionally, in IoT settings, the application of a combination of deep learning approaches and

interpretability alongside language-based explanations is proven to increase the efficiency of detections and improve transparency (Alasmari et al., 2025). Explainability is especially important for addressing issues of authenticity and authorship associated with content generation. The growing complexity of the methods used for generating text creates the problem of how one can distinguish human-generated text from the machine-generated one. Several studies have shown that the use of machine learning approaches along with explainability allows classifying and attributing text properly while determining distinctive linguistic and structural features of the text (Najjar et al., 2025).

Beyond its uses in various fields, explainability must also take into consideration the human perspective. Many researchers have found that transparency does not only comprise technical features but also encompasses a wide array of other considerations. Human perspectives highlight the significance of including users, ethics and participatory design in developing interpretable machines (Ehsan et al., 2025). The importance of explainability comes out very clearly when it is applied in collaborative machines. Some recent researches have proposed architectural approaches to incorporate predictive modeling, explanations and user interfaces in ways that support the effective collaboration between human and machine (Jeong, 2025). Actually, architectural approaches have proved their effectiveness in helping users make decisions in complicated situations from medical diagnosis to strategic planning.

Lastly, explainability has been used in the sphere of cultural heritage and archives and contributed to improving accessibility and interpretations of digitized documents. Thanks to the application of computational methods along with visualization and explanation technologies, scientists have managed to show that explainability may help to render historical information more accessible to the general audience (Toth et al., 2025). Nevertheless, a number of issues concerning explainability exist. For example, one of the issues connected with XAI is the question of finding a balance between simplification and fidelity because overly simple explanations will not adequately capture the phenomenon, while fidelity explanations will be very complex and hard to understand (Arous et al., 2025). Moreover, such issues as the generation of explanations that are entirely incorrect, absence of standards regarding evaluations, inconsistency of explanations are the problems standing in the way of developing XAI (Ehsan et al., 2025; Pehlke & Jansen, 2025).

Table.4. Explainability Score Comparison (5-point Scale)

Method	Score
LLM	2.1
BERT	3.2
SHAP	4.8
LIME	4.6
SHAP + LLM	4.9
Hybrid Explainable Models	5.0

An analysis of explainability from Table 4 reveals substantial disparities among the examined approaches. Standalone large language models earn the lowest marks as decision-making process is extremely complex and difficult to comprehend in them. Systems using BERT perform slightly better; however, additional mechanisms are required in order to provide greater clarity of the decision-making process. Techniques such as SHAP and LIME achieve far better results due to the fact that they provide clear insight into the influence of features on predictions. The highest explainability is provided by hybrid frameworks as they demonstrate high predictive performance and clarity of explanations. It proves that having built-in explainability is crucial in increasing the level of trust, compliance and implementation of AI models in practice. The potential solution to the problem consists in the development of modular and auditable pipelines that allow one to analyze the whole process of reasoning performed by the machine and validate its outcomes (Pehlke & Jansen, 2025). Besides, interdisciplinary collaboration is gaining more importance as an essential element of efficient application of the technology that requires significant efforts in order to be interpreted.

III.Survey on Existing Attention models applied on Large Language Models

Such developments have transformed the whole field of artificial intelligence in terms of creating highly advanced multimodal understanding, contextual reasoning and intelligent decision making. Introduction of the transformer architecture and self-attention mechanism resulted in increased ability of modern models to deal with large-scale sequential and multimodal data sets. Currently, the research concentrates on the development of explainability of the models and reduction of the problem of hallucination along with optimization of domain adaptation and external knowledge representations for better performance. This section will present the recent results on attention mechanisms, multimodal models, knowledge-aware architectures, explainable AI and domain-specific LLMs.

The transformer architecture which is based on the use of self-attention mechanism is the core of currently developed large language and vision-language models. Luo et al. (2023) claim that self-attention mechanism allows the model to efficiently capture long-range context and dependencies relations. As a result, transformers have been transformed from the models used in machine translation tasks into the core models used in generative AI models. The main feature of the transformers when it comes to handling sequences is their ability to parallelize the process. In Li et al.'s (2024) paper, there is an extensive review of efficient LLM training and inference approaches. The researchers stressed the computational costliness of large models training in institutions with a small number of GPUs. The authors discussed such optimization approaches as parameter-efficient tuning, distributed training, quantization and memory-efficient inference, proving that the usage of billion-parameter models with limited infrastructure is possible.

Moreover, the application of attention-based mechanisms has been further expanded to multimodal systems. Yu et al. (2024) developed a multimodal prompting approach called Attention Prompting on Image (API). It incorporates the use of text-query-guided attention heatmaps in conjunction with images. The experiments have shown significant improvement on the vision-language benchmarks, suggesting that attention-guided interaction works well for multimodal tasks. Similarly, Liu et al. (2024) solved the problem of hallucinations in LVLMs with a training-free attention modification technique. In their paper, the idea of "text inertia" was stated – according to which LVLMs generate outputs that are strongly affected by text and not by images. Increasing attention weights for image tokens and reducing those for textual reasoning allowed significantly decreasing hallucinations in various LVLMs.

The use of knowledge bases within LLMs is a growing trend towards building more fact-aware and reasoning-savvy language models. For instance, Yang et al. (2024) conducted a study about the role of knowledge graphs in enhancing fact-awareness of large language models through knowledge-driven text generation. To this end, the authors proposed a novel architecture called KGLLM which combines semantics and probabilistic language modeling. In the medical field, Li et al. (2024) developed the LLM-based Knowledge-Aware Attention Network (LKAN) for liver cancer staging using radiology reports. In this regard, the LKAN system combines domain-specific pre-training and attention-based feature selection with a global-local attention-based mechanism, allowing the extraction of staging-related features from radiology reports. Overall, the model performed better than other baseline methods.

The field of mental health care has also benefitted from the use of customized LLM architectures. Hu et al. (2024) presented PsychoLLM, which was a psychological language model designed using high-quality datasets involving question-answer sessions, dialogues and psychological knowledge exchange processes. This research study stressed the importance of domain-specific fine-tuning and evidenced-based conversation for improved psychological reasoning and evaluations. Nikbakht et al. (2024) developed TSpec-LLM for the telecommunication sector; TSpec-LLM was an open-source dataset comprised of 3GPP specifications from Release 8 to Release 19. They utilized LLM along with retrieval-augmented generation method to enhance question and answer performance concerning technical information on telecommunications literature. They observed that the integration of contextualized retrieval substantially increased the precision rates of LLM. Last, Jiang et al. (2024) invented Eplus-LLM, which was an automatic energy building modeling tool used for transforming descriptions into EnergyPlus models. Fine-tuning the model was vital for the enhancement of engineering modeling processes.

Multimodal learning innovations have become common in recent advancements in AI systems. For example, Song et al. (2024) innovated MoMA, a new multimodal LLM adapter that can generate personalized images according to users' preferences. Through combining the use of image and text

data through a two-stage feature extraction and generation approach, MoMA was able to generate identity-preserved images while retaining faithfulness to the prompts used. The proposed self-attention shortcut technique, for instance, allowed the authors to successfully encode visual data into the generative models. Other use cases of the multimodal LLM architectural framework involved the application of the architecture to the task of trajectory prediction for intelligent systems. To this effect, Lan et al. (2024) proposed a new multimodal architecture, Traj-LLM for autonomous vehicles. Traj-LLM applies sparse context encoding, lane-attentive probabilistic learning, and multimodal decoding approaches for better traffic trajectory predictions. From the experiments performed, it is clear that the capacity of the pretrained LLMs to understand semantics greatly improves the scene understanding. In the field of keyword extraction, many studies applied the power of LLMs to improve the performance of the task. In this regard, Maragheh et al. (2023) came up with a multimodal approach to keyword extraction known as LLM-TAKE. LLM-TAKE uses contextual inference and hallucination reduction techniques to generate extractive and abstractive thematic keywords.

As the number of uses of LLMs grows, explainability and transparency are becoming more and more relevant. According to Haurigné et al. (2024), there is a vulnerability detection method that exploits the capabilities of LLMs based on BERT and explainable AI methods, namely SHAP, LIME, and attention heat maps. This model was introduced as the solution to one of the limitations of black box approaches to deep learning models in that it provided transparency by determining the most influential tokens that led to the prediction of the presence of vulnerabilities. Also, there are applications of LLMs to cybersecurity problems. For instance, Otal and Canbaz (2024) developed an interactive honeypot that used LLMs. Using fine-tuned pretrained language models and datasets of attacker generated commands, this model could simulate realistic interactions with the attacker and generate context-aware outputs to help analyze the threat and detect any malicious activities. Nevertheless, there are still some improvements in reliability and ethical aspects that can be made. As LLMs start being used in healthcare, these become highly relevant. It can be shown by the study done by Ullah et al. (2024) on the use of LLMs in medicine and digital pathology.

Table.5. Computational Complexity Analysis

Model	Training Cost	Inference Cost
Traditional ML	Low	Low
BERT	Medium	Medium
RoBERTa	Medium	Medium
GPT-2	High	High
LLMs	Very High	Very High
Hybrid Models	Very High	High

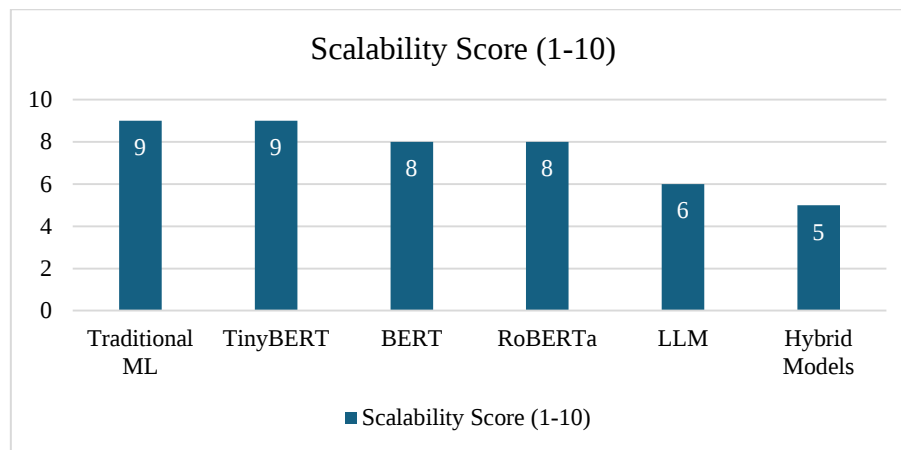
According to Table 5, the more complex a model is, the more resources it needs. Traditional machine learning methods have the lowest costs for both training and inference, making them perfect for real-time and resource-limited use. BERT-based systems need moderate resources, giving good performance without being too costly. Larger language models, however, use a lot of computing power because of how big they are and how much work goes into training them. Hybrid architectures also come with high costs due to the complexity of putting different components together. These findings show that while more advanced architectures usually give better predictions, deploying them can be tough because of limited hardware, energy use and running costs.

Table.6. Scalability Evaluation

Approach	Scalability Score (1-10)
Traditional ML	9
TinyBERT	9
BERT	8
RoBERTa	8
LLM	6
Hybrid Models	5

The scalability evaluation shows that traditional machine learning models and lighter transformer versions top the scalability charts. With simpler structures and less resource needs, they work well in various settings. Plus, Fig.2., illustrates this.

Fig.2. Scalability Evaluation



BERT and RoBERTa also do well, especially on specialized tasks. However, bigger language models get lower scores due to high memory use and compute needs. Furthermore, hybrid designs fare worst, since combining different parts makes implementation and upkeep harder. All in all, these results indicate scalability continues to be a big hurdle for large smart systems, despite their great performance traits.

Table.7. Generalization Capability

Model	Cross-Domain Score
Traditional ML	5.2
BERT	7.4
RoBERTa	7.8
GPT Models	9.1
Fine-Tuned LLMs	9.4
Hybrid Models	8.8

From Table 7, it can be seen that fine-tuned large language models show superiority in transferring knowledge between multiple domains. The reason for such performance is the large amount of pre-training of such models, during which they manage to capture multiple language regularities and contexts and become more flexible in performing tasks related to different purposes. At the same time, models based on GPT demonstrate high results in terms of cross-domain generalization. Hybrid approaches continue to show high performance due to utilizing advantages of different learning approaches. However, BERT-based models generally require additional task-specific fine-tuning in order to perform perfectly in different domains, even though they still demonstrate high performance. Meanwhile, traditional machine learning approaches demonstrate poor flexibility due to their reliance on handcrafted features and particular data sets. Consequently, the study proves that modern language models became much more effective in working with multiple domains.

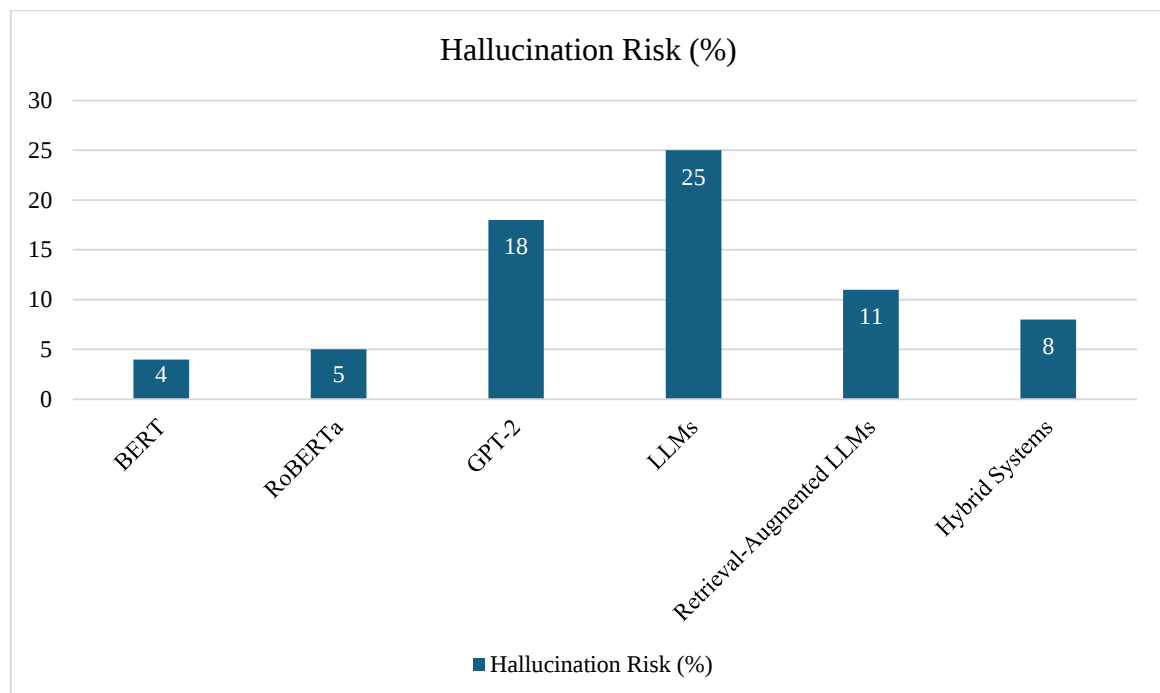
Table.8. Hallucination Risk Assessment

Model	Hallucination Risk (%)
BERT	4
RoBERTa	5

GPT-2	18
LLMs	25
Retrieval-Augmented LLMs	11
Hybrid Systems	8

Table 8 shows that large language models have the highest hallucination risk. This happens because they can spout fluent but inaccurate info when uncertain. GPT-based models suffer from hallucinations too, though. On the bright side, retrieval-augmented systems fare better since they tie their output to external info, boosting accuracy. Similarly, hybrid architectures cut down on hallucinations by adding validation processes. BERT and RoBERTa models are least prone to this problem; they focus on tasks like rep learning and classification not open-ended text creation.

Fig.3. Hallucination Risk Assessment



These results highlight the need to integrate checks and external knowledge in systems that generate text.

Embeddings obtained using LLMs have proven to be very effective for representing text data and clustering. Keraghel et al. (2024) compared various LLM embeddings used for document clustering. The experiment revealed the efficiency of embeddings created by large transformers as opposed to conventional methods. However, it also proved that there were some drawbacks in terms of computational power when using larger embedding models. The capacity of the transformer embeddings to incorporate contextual semantics has made them more versatile for various NLP applications such as classification, information retrieval, summarization and semantic search.

IV.Study on Attention Models to enhance LLMs

Large language models have experienced fast growth through advances in the design of transformers, where the attention module acts as the core that supports context and sequence learning. Attention enables models to concentrate on essential data segments, improving their performance over traditional sequential processes. Consequently, this innovation has advanced models' efficiency and the capability to learn from massive data and capture long-distance dependencies (Liu, 2025). It has enabled models to understand the relationships among words and phrases by analyzing them based on context and semantics. Thus, attention mechanisms have been instrumental in achieving reasoning, contextual alignment and long text comprehension.

Attention heads are the key elements in transformers, which provide particular linguistic and semantic connections between words and phrases. Attention heads operate in the same way as human cognitive abilities do when performing reasoning and understanding tasks (Zheng et al., 2025). The exact work of attention heads in transformers is still unknown because most deep-learning algorithms fail to explain the process of their functioning. The behaviour of attention heads has been described and interpreted, giving rise to models connecting their functions to particular phases of human reasoning. Along with being more interpretable, attention mechanisms have been found to increase memory retrieval and contextual searches. New research in generative agents shows how cross-attention networks can be used for better memory retrieval through dynamic prioritization of the importance of the information in comparison with the input data (Hong & He, 2025). Thus, the implementation of the attention-based algorithms in memory systems would greatly contribute to the capabilities of intelligent agents as well as their ability to adapt to the environment. However, the issues of efficiency and scalability become more relevant as the amount of processed data increases. Several scientists have proposed some approaches for dealing with these problems. Optimization of prompts based on attention has proved to be efficient in lowering the complexity of the attention mechanism, as it allows compressing prompts without lowering their quality (Zhao et al., 2025). Along with prompt optimization, attention mechanisms can be improved through pruning of unnecessary features to decrease the computation time (Belhaouari & Kraidia, 2025). In terms of speeding up the inference process, scientists suggest using clustered attention mechanisms and hierarchical key selection to reduce computational complexity and work only with essential data elements.

However, while improving model performance, attention mechanisms increase its vulnerability to certain risks. Based on the research into adversarial attacks, attention mechanisms can become an object of prompt injection attacks and multi-turn interactions (Hung et al., 2025; Du et al., 2025). The attention-based approach was proven to be useful for detecting these threats. Hence, a deep knowledge of the working principles of attention systems is needed for providing safety of the systems. The hybrid architecture with the use of recurrent networks in combination with the attention layers allows distinguishing between different classes of texts created by both computers and people. The focus on the required linguistic properties of the texts makes it possible to ensure the integrity of the information and overcome any problems with the automatic generation of texts. In addition to developments within the framework of algorithms, the attention-based architecture is actively applied in the specialized tasks as well. For instance, the agricultural tasks which involve the use of large attention-based systems in making decisions and improving the process are effective even though they lag behind in innovations in comparison with other spheres (Shaikh et al., 2025).

Table.9. Attention Mechanism Impact

Attention Technique	Improvement (%)
Self-Attention	12.3
Multi-Head Attention	15.4
Cross Attention	18.7
Knowledge-Aware Attention	21.6
Multimodal Attention	24.3

As Table.9. suggests, advanced attention mechanisms greatly improve model performance in different tasks. Basic self-attention provides a significant advantage because of its ability to detect input dependencies. Multi-head attention adds more value due to its ability to identify multiple context dependencies. Cross-attention increases effectiveness because it allows inputs to communicate with each other. Knowledge-aware attention works even better as it uses structured external information. The highest improvement is achieved by multimodal attention because it combines textual, visual, and sensor input. All of these examples demonstrate the importance of attention mechanisms for developing intelligent models. There are new trends that appear about combining structured knowledge and human behavior models into attention mechanisms. Attention mechanisms based on the injection of knowledge allow for an effective extraction and updating of information inside the model, which helps the model perform well in knowledge-intensive tasks (Yu et al., 2026).

Additionally, the use of human attention such as eye tracking information has proven to be effective in increasing the focus of models while performing code summarization tasks (Zhang et al., 2026).

There are numerous challenges to overcome, however. Although attention mechanisms provide many opportunities, there are concerns that they may increase the number of issues related to efficiency, bias and unintended behavior. The challenge of striking the right balance between performance and interpretability should also be addressed because the more complex an attention mechanism is, the less transparent and explainable it becomes. Moreover, there is the question of ethics in decision-making and the use of technologies by humans. In this regard, the problem of moral reasoning is discussed (Graves, 2025). All in all, it can be argued that attention mechanisms play an important role in current language-based systems as they ensure high-quality performance and scalability of the models. Research is being conducted on further enhancement of attention mechanisms by improving efficiency, understanding and security of these processes and incorporating human knowledge into them. Overcoming present shortcomings is an important condition for proper development of these systems.

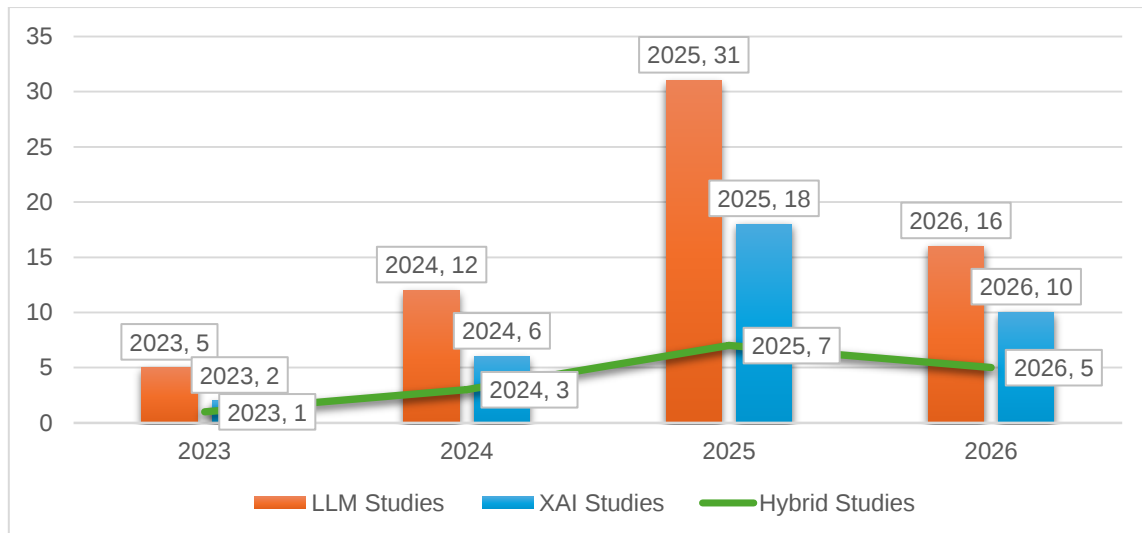
For future research, it would be beneficial to develop symbolic reasoning in neural networks to increase consistency and reduce hallucination effects (Papagiannopoulos et al., 2025). The energy efficiency via model compression and architectural optimization should be considered (Rahulamathavan et al., 2025). Adaptive models capable of processing dynamic streams of information should be developed for advancing cybersecurity and IoT (Rahman et al., 2025). The use of multiple types of data, including texts, images, and sensors may bring many benefits for healthcare and intelligent systems (Annepaka & Pakray, 2025). It is important to guarantee fairness, accountability, and transparency of models by introducing appropriate measures for model evaluation and regulating policy (Yin et al., 2025). In general, based on the above comparison, it can be concluded that BERT-like models and classic models continue dominating in structured environment whereas LLMs prevail in semantic tasks.

Table.19. Trend Analysis (2023–2026)

Year	LLM Studies	XAI Studies	Hybrid Studies
2023	5	2	1
2024	12	6	3
2025	31	18	7
2026	16	10	5

Table.19., shows how research into Large Language Models (LLMs), Explainable Artificial Intelligence (XAI) and mixes of both has grown from 2023 to 2026. There's been a big rise in work across all three areas. LLM studies shot up from 5 papers in 2023 to 31 in 2025. This highlights the fast-growing interest in advanced language tech. At the same time, explainability research also boomed—from 2 papers in 2023 to 18 in 2025. That shows more scientists care about transparency, accountability and understanding these models better.

Fig.6. Trend Analysis of the research works (2023-2026)



Hybrid projects saw growth too, moving from one study in 2023 to seven in 2025. However, there seems to be fewer studies in 2026. Yet, this drop might just mean some late-year research wasn't included, not that people aren't interested anymore. In the end, the numbers point to a move towards creating AI that's easier to understand while still performing well. Researchers want models that work great but make sense to use in real life situations as well.

Table.20. Performance Ranking of Approaches

Rank	Approach
1	Hybrid Models
2	Fine-Tuned LLMs
3	RoBERTa
4	BERT
5	Traditional ML
6	GPT-2

In Table.20. one can clearly see which computational techniques perform best in comparison with others in terms of accuracy, explainability, generalization, scalability, and efficiency. The leading position belongs to hybrid models because of their ability to combine strengths of various architectural paradigms while reducing their weaknesses. The second position belongs to fine-tuned Large Language Models which are performing incredibly well in prediction tasks and across various domains. RoBERTa occupies the third place and BERT the fourth one; both of them have demonstrated their capability to deal with structured language processing tasks. Traditional machine learning approaches take the fifth place being not as accurate as hybrid models, but being more efficient and scalable. Finally, GPT-2 takes the last position being beaten by other models regarding accuracy, adaptability and explainability. One can conclude from this table that integrated architectures which are characterized by balanced performance and practicality become extremely important for solving real-world tasks.

V.Conclusion

This review covered the evolution, applications, and limitations of explainability and reliability of large language model systems in healthcare, education, cybersecurity, finance, industrial applications and IoT. The information presented in this review shows that the performance of the system is not enough to be used in the decision-making process when transparency, accountability, and trustworthiness are critical. In this context, SHAP and LIME offer better explanations of feature importance compared to other approaches. At the same time, the combination of those techniques with the language models, including hybrid frameworks, can make model outputs more understandable.

Based on the comparison presented in this review, the hybrid explainable models had the highest explainability score, followed by the SHAP with LLMs, LIME, SHAP, BERT and LLMs.

However, explainability is still hindered by the trade-off between fidelity and explainability. The problems associated with hallucinations, bias, privacy issues, resource usage, inconsistency of explanations and the lack of evaluation criteria present an obstacle to the reliable application of the technique. Thus, the key requirement for the explanation methods can be considered as their fidelity, efficiency, domain specificity and suitability.

References

- [1] Ahmed, A., Saleem, M., Alzeen, M., Birur, B., Fargason, R. E., Burk, B. G., ... & Al-Garadi, M. A. (2025). Explainable AI for mental health emergency returns: integrating LLMs with predictive modeling. *arXiv preprint arXiv:2502.00025*. <https://doi.org/10.48550/arXiv.2502.00025>
- [2] Alasmari, S. M., Sakly, H., Kraiem, N., & Algarni, A. (2025). Phishing detection in IoT: an integrated CNN-LSTM framework with explainable AI and LLM-enhanced analysis. *Discover Internet of Things*, 5(1), 102. <https://doi.org/10.1007/s43926-025-00202-9>
- [3] Annepaka, Y., & Pakray, P. (2025). Large language models: a survey of their development, capabilities and applications. *Knowledge and Information Systems*, 67(3), 2967-3022. <https://doi.org/10.1007/s10115-024-02310-4>
- [4] Arous, I., Chehbouni, K., Cheng, Z., & Dossou, B. (2025). Llm explainability. In *Handbook of Human-Centered Artificial Intelligence* (pp. 1-61). Singapore: Springer Nature Singapore. https://doi.org/10.1007/978-981-97-8440-0_85-1
- [5] Arslan, B., Nuhoglu, C., Satıcı, M. O., & Altınbilek, E. (2025). Evaluating LLM-based generative AI tools in emergency triage: a comparative study of ChatGPT Plus, Copilot Pro and triage nurses. *The American journal of emergency medicine*, 89, 174-181. <https://doi.org/10.1016/j.ajem.2024.12.024>
- [6] Belhaouari, S. B., & Kraidia, I. (2025). Efficient self-attention with smart pruning for sustainable large language models. *Scientific Reports*, 15(1), 10171. <https://doi.org/10.1038/s41598-025-92586-5>
- [7] Bilal, A., Ebert, D., & Lin, B. (2025). LLMs for explainable ai: A comprehensive survey. *arXiv preprint arXiv:2504.00125*. <https://doi.org/10.48550/arXiv.2504.00125>
- [8] Chakraborty, C., Bhattacharya, M., Pal, S., Chatterjee, S., Das, A., & Lee, S. S. (2025). AI-enabled language models (LMs) to large language models (LLMs) and multimodal large language models (MLLMs) in drug discovery and development. *Journal of advanced research*. <https://doi.org/10.1016/j.jare.2025.02.011>
- [9] Chittimalla, S. K., & Potluri, L. K. M. (2025, March). Explainable AI frameworks for large language models in high-stakes decision-making. In *2025 International Conference on Advanced Computing Technologies (ICoACT)* (pp. 1-6). IEEE. DOI: [10.1109/ICoACT63339.2025.11005089](https://doi.org/10.1109/ICoACT63339.2025.11005089)
- [10] Doshi, A. R., Bell, J. J., Mirzayev, E., & Vanneste, B. S. (2025). Generative artificial intelligence and evaluating strategic decisions. *Strategic Management Journal*, 46(3), 583-610. <https://doi.org/10.1002/smj.3677> **Digital Object Identifier (DOI)**
- [11] Du, X., Mo, F., Wen, M., Gu, T., Zheng, H., Jin, H., & Shi, J. (2025, April). Multi-turn jailbreaking large language models via attention shifting. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 39, No. 22, pp. 23814-23822). <https://doi.org/10.1609/aaai.v39i22.34553>
- [12] Ehsan, U., Watkins, E. A., Wintersberger, P., Manger, C., Hubig, N., Savage, S., ... & Riener, A. (2025, April). New frontiers of human-centered explainable AI (HCXAI): Participatory civic AI, benchmarking LLMs, XAI hallucinations and responsible AI audits. In *Proceedings of the Extended Abstracts of the CHI Conference on Human Factors in Computing Systems* (pp. 1-6). <https://doi.org/10.1145/3706599.3706713>
- [13] Feldman, M. J., Hoffer, E. P., Conley, J. J., Chang, J., Chung, J. A., Jernigan, M. C., ... & Chueh, H. C. (2025). Dedicated AI expert system vs generative AI with large language model for clinical diagnoses. *JAMA Network Open*, 8(5), e2512994. [doi:10.1001/jamanetworkopen.2025.12994](https://doi.org/10.1001/jamanetworkopen.2025.12994)
- [14] Ferrag, M. A., Alwahedi, F., Battah, A., Cherif, B., Mechri, A., Tihanyi, N., ... & Debbah, M. (2025). Generative AI in cybersecurity: A comprehensive review of LLM applications and vulnerabilities. *Internet of Things and Cyber-Physical Systems*, 5, 1-46. <https://doi.org/10.1016/j.iotcps.2025.01.001>
- [15] Graves, M. (2025). Moral attention is all you need. *Theology and Science*, 23(2), 241-248. <https://doi.org/10.1080/14746700.2025.2472118>

- [16] Haurogné, J., Basheer, N., &Islam, S. (2024). Vulnerability detection using BERT based LLM model with transparency obligation practice towards trustworthy AI. *Machine Learning with Applications*, 18, 100598. <https://doi.org/10.1016/j.mlwa.2024.100598>
- [17] Hong, C., &He, Q. (2025). Enhancing memory retrieval in generative agents through LLM-trained cross attention networks. *Frontiers in Psychology*, 16, 1591618. <https://doi.org/10.3389/fpsyg.2025.1591618>
- [18] Hu, J., Dong, T., Luo, G., Ma, H., Zou, P., Sun, X., ... &Wang, M. (2024). Psycollm: Enhancing llm for psychological understanding andevaluation. *IEEE Transactions on Computational Social Systems*, 12(2), 539-551. DOI: [10.1109/TCSS.2024.3497725](https://doi.org/10.1109/TCSS.2024.3497725)
- [19] Hung, K. H., Ko, C. Y., Rawat, A., Chung, I. H., Hsu, W. H., &Chen, P. Y.(2025, April). Attention tracker: Detecting prompt injection attacks in llms. In *Findings of the Association for Computational Linguistics: NAACL 2025* (pp. 2309-2322). [10.18653/v1/2025.findings-naacl.123](https://doi.org/10.18653/v1/2025.findings-naacl.123)
- [20] Jauhiaiainen, J. S., &Garagorry Guerra, A. (2025). Generative AI in education: ChatGPT-4 in evaluating students' written responses. *Innovations in Education and Teaching International*, 62(4), 1377-1394. <https://doi.org/10.1080/14703297.2024.2422337>
- [21] Jeong, C. (2025). Design and evaluation methods forLLM-based explainable AI (XAI)-based human-AI collaboration systems. *Advances in Artificial Intelligence and Machine Learning*, 5(3), 240.
- [22] Jiang, G., Ma, Z., Zhang, L., &Chen, J. (2024). EPlus-LLM: A large language model-based computing platform for automated building energymodeling. *Applied Energy*, 367, 123431. <https://doi.org/10.1016/j.apenergy.2024.123431>
- [23] Jiao, J., Afroogh, S., Chen, K., Atkinson, D., &Dhurandhar, A. (2025). Generative AI and LLMs in industry: A text-mining analysis and critical evaluation of guidelines and policy statements across fourteen industrialsectors. *arXiv preprint arXiv:2501.00957*. <https://doi.org/10.48550/arXiv.2501.00957>
- [24] Kadoma, K., Metaxa, D., &Naaman, M. (2025, April). Generative AI and Perceptual Harms: Who's Suspected of using LLMs?. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems* (pp. 1-17). <https://doi.org/10.1145/3706598.3713897>
- [25] Keraghel, I., Morbieu, S., &Nadif, M. (2024, April). Beyond words: a comparative analysis of LLM embeddings for effective clustering. In *International symposium on intelligent data analysis* (pp. 205-216). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-58547-0_17
- [26] Kumar, S., Praveen, R. V. S., Aida, R., Varshney, N., Alsalami, Z., &Boob, N. S. (2025, September). Enhancing AI Decision-Making with Explainable Large Language Models (LLMs) in Critical Applications. In *2025 IEEE International Conference on Advances in Computing Research On Science Engineering and Technology (ACROSET)* (pp. 1-6). IEEE. DOI: [10.1109/ACROSET66531.2025.11280656](https://doi.org/10.1109/ACROSET66531.2025.11280656)
- [27] Lan, Z., Liu, L., Fan, B., Lv, Y., Ren, Y., &Cui, Z. (2024). Traj-llm: A new exploration for empowering trajectory prediction with pre-trained large language models. *IEEE Transactions on Intelligent Vehicles*. DOI: [10.1109/TIV.2024.3418522](https://doi.org/10.1109/TIV.2024.3418522)
- [28] Lee, R. (2025). Large Language Models (LLMs) and Generative Artificial Intelligence (GenAI). In *Natural Language Processing: A Textbook with Python Implementation* (pp. 241-273). Singapore: Springer Nature Singapore. https://doi.org/10.1007/978-981-96-3208-4_10
- [29] Li, R., Fu, D., Shi, C., Huang, Z., &Lu, G. (2024). Efficient LLMs training and inference: An introduction. *IEEE Access*, 13, 32944-32970. DOI: [10.1109/ACCESS.2024.3501358](https://doi.org/10.1109/ACCESS.2024.3501358)
- [30] Li, Y., Zheng, X., Li, J., Dai, Q., Wang, C. D., &Chen, M. (2024). LKAN: LLM-based knowledge-aware attention network for clinical staging of liver cancer. *IEEE Journal of Biomedical and Health Informatics*, 29(4), 3007-3020. DOI: [10.1109/JBHI.2024.3478809](https://doi.org/10.1109/JBHI.2024.3478809)
- [31] Lim, B., Huerta, R., Sotelo, A., Quintela, A., &Kumar, P. (2025). Explicate: Enhancing phishing detection through explainable ai and llm-powered interpretability. *arXiv preprint arXiv:2503.20796*. <https://doi.org/10.48550/arXiv.2503.20796>
- [32] Liu, J., Sihes, A. J. B., &Lu, Y. (2025). How do generative artificial intelligence (AI) tools and large language models (LLMs) influence language learners' critical thinking in EFL education? A systematic review. *Smart Learning Environments*, 12(1), 48. <https://doi.org/10.1186/s40561-025-00406-0>
- [33] Liu, S., Zheng, K., &Chen, W. (2024, September). Paying more attention to image: A training-free method for alleviating hallucination in llms. In *European Conference on Computer Vision* (pp. 125-140). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-73010-8_8

- [34] Liu, Y.(2025). Attention is all large language model need. In *ITM Web of Conferences* (Vol. 73, p. 02025). EDP Sciences. <https://doi.org/10.1051/itmconf/20257302025>
- [35] López-Pernas, S., Song, Y., Oliveira, E., & Saqr, M.(2025). LLMs for explainable artificial intelligence: Automating natural language explanations of predictive analytics models. In *Advanced Learning Analytics Methods: AI, Precision and Complexity* (pp. 261-286). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-95365-1_11
- [36] Luo, Q., Zeng, W., Chen, M., Peng, G., Yuan, X., & Yin, Q. (2023, July). Self-attention and transformers: Driving the evolution of large language models. In *2023 IEEE 6th International conference on electronic information and communication technology (ICEICT)* (pp. 401-405). IEEE. DOI: [10.1109/ICEICT57916.2023.10245906](https://doi.org/10.1109/ICEICT57916.2023.10245906)
- [37] Maragheh, R. Y., Fang, C., Irugu, C. C., Parikh, P., Cho, J., Xu, J., ... & Achan, K. (2023, December). LLM-TAKE: Theme-aware keyword extraction using large language models. In *2023 IEEE International Conference on Big Data (BigData)* (pp. 4318-4324). IEEE. DOI: [10.1109/BigData59044.2023.10386476](https://doi.org/10.1109/BigData59044.2023.10386476)
- [38] McIntosh, T. R., Susnjak, T., Arachchilage, N., Liu, T., Xu, D., Watters, P., & Halgamuge, M. N. (2025). Inadequacies of large language model benchmarks in the era of generative artificial intelligence. *IEEE Transactions on Artificial Intelligence*. DOI: [10.1109/TAI.2025.3569516](https://doi.org/10.1109/TAI.2025.3569516)
- [39] Mumuni, F., & Mumuni, A.(2025). Explainable artificial intelligence (XAI): from inherent explainability to large language models. *arXiv preprint arXiv:2501.09967*. <https://doi.org/10.48550/arXiv.2501.09967>
- [40] Najjar, A. A., Ashqar, H. I., Darwish, O., & Hammad, E. (2025). Leveraging explainable ai for llm text attribution: Differentiating human-written and multiple llm-generated text. *Information*, 16(9), 767. <https://doi.org/10.3390/info16090767>
- [41] Nezami, Z., Hafeez, M., Djemame, K., & Zaidi, S. A. R. (2025, June). Generative AI on the edge: Architecture and performance evaluation. In *ICC 2025-IEEE International Conference on Communications* (pp. 4595-4602). IEEE. DOI: [10.1109/ICC52391.2025.11161569](https://doi.org/10.1109/ICC52391.2025.11161569)
- [42] Nikbakht, R., Benzaghta, M., & Geraci, G. (2024, December). Tspec-llm: An open-source dataset for llm understanding of 3gpp specifications. In *2024 IEEE Globecom Workshops (GC Wkshps)* (pp. 1-6). IEEE. DOI: [10.1109/GCWkshp64532.2024.11101012](https://doi.org/10.1109/GCWkshp64532.2024.11101012)
- [43] Ong, J. C. L., Chen, M. H., Ng, N., Elangovan, K., Tan, N. Y. T., Jin, L., ... & Liu, N. (2025). A scoping review on generative AI and large language models in mitigating medication related harm. *npj Digital Medicine*, 8(1), 182. <https://doi.org/10.1038/s41746-025-01565-7>
- [44] Otal, H. T., & Canbaz, M. A. (2024, September). Llm honeypot: Leveraging large language models as advanced interactive honeypot systems. In *2024 IEEE Conference on Communications and Network Security (CNS)* (pp. 1-6). IEEE. DOI: [10.1109/CNS62487.2024.10735607](https://doi.org/10.1109/CNS62487.2024.10735607)
- [45] Papagiannopoulos, I., Koutalidis, H., Rempi, P., Ntanos, C., & Askounis, D. (2025). Comparison of explainability methods for hallucination analysis in LLMs. *Open Research Europe*, 5, 191. <https://doi.org/10.12688/openreseurope.20839.1>
- [46] Pehlke, M., & Jansen, M. (2025). Llm driven processes to foster explainable ai. *arXiv preprint arXiv:2511.07086*. <https://doi.org/10.48550/arXiv.2511.07086>
- [47] Rahman, M. A., Francia, G., Shahriar, H., Cuzzocrea, A., Mohamed, A., Khan, M. U., & Ahamed, S. I. (2025, December). A Survey of Large Language Models (LLMs) for Cybersecurity: Opportunities and Directions. In *2025 IEEE International Conference on Big Data (BigData)* (pp. 4333-4342). IEEE. DOI: [10.1109/BigData66926.2025.11402639](https://doi.org/10.1109/BigData66926.2025.11402639)
- [48] Rahulamathavan, Y., Farooq, M., & De Silva, V. (2025). PLEX: Perturbation-free local explanations for llm-based text classification. *IEEE Journal of Selected Topics in Signal Processing*, 19(7), 1266-1278. DOI: [10.1109/JSTSP.2025.3633593](https://doi.org/10.1109/JSTSP.2025.3633593)
- [49] Rouzrokh, P., Khosravi, B., Faghani, S., Moassefi, M., Shariatnia, M. M., Rouzrokh, P., & Erickson, B. (2025). A current review of generative AI in medicine: core concepts, applications and current limitations. *Current Reviews in Musculoskeletal Medicine*, 18(7), 246-266. <https://doi.org/10.1007/s12178-025-09961-y>
- [50] Shaikh, T. A., Rasool, T., Venington, K., & Yaseen, S. M. (2025). The role of large language models in agriculture: harvesting the future with LLM intelligence. *Progress in Artificial Intelligence*, 14(2), 117-164. <https://doi.org/10.1007/s13748-024-00359-4>

- [51] Song, K., Zhu, Y., Liu, B., Yan, Q., Elgammal, A., &Yang, X. (2024, September). Moma: Multimodal llm adapter for fast personalized image generation. In *European Conference on Computer Vision* (pp. 117-132). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-73661-2_7
- [52] Toth, G. M., Albrecht, R., &Pruski, C.(2025). Explainable AI, LLM anddigitized archival cultural heritage: A case study of the Grand Ducal Archive of the Medici. *AI & SOCIETY*, 40(6), 4561-4573. <https://doi.org/10.1007/s00146-025-02238-5>
- [53] Ullah, E., Parwani, A., Baig, M. M., &Singh, R. (2024). Challenges and barriers of using large language models (LLM) such as ChatGPT for diagnostic medicine with a focus on digital pathology– a recent scoping review. *Diagnostic pathology*, 19(1), 43. <https://doi.org/10.1186/s13000-024-01464-7>
- [54] Xiong, L., Wang, H., Chen, X., Sheng, L., Xiong, Y., Liu, J., ... &Tang, Y. (2025). DeepSeek: paradigm shifts and technical evolution in large AI models. *IEEE/CAA Journal of Automatica Sinica*, 12(5), 841-858. DOI: [10.1109/JAS.2025.125495](https://doi.org/10.1109/JAS.2025.125495)
- [55] Yang, L., Chen, H., Li, Z., Ding, X., &Wu, X. (2024). Give us the facts: Enhancing large language models with knowledge graphs for fact-aware language modeling. *IEEETransactions on Knowledge and Data Engineering*, 36(7), 3091-3110. DOI: [10.1109/TKDE.2024.3360454](https://doi.org/10.1109/TKDE.2024.3360454)
- [56] Yigit, Y., Ferrag, M. A., Ghanem, M. C., Sarker, I. H., Maglaras, L. A., Chrysoulas, C., ... &Janicke, H. (2025). Generative AI and LLMs for critical infrastructure protection: evaluation benchmarks, agentic AI, challenges andopportunities. *Sensors*, 25(6), 1666. <https://doi.org/10.3390/s25061666>
- [57] Yin, F., Zhong, M.,&Ru, Z. (2025). Exploring explainability in large language models. doi: 10.20944/preprints202503.2318.v1
- [58] Yu, B., Huang, W., &Liu, K. (2026, March). SR-KI: Scalable and Real-Time Knowledge Integration into LLMs via Supervised Attention. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 40, No. 41, pp. 34486-34494). <https://doi.org/10.1609/aaai.v40i41.40747>
- [59] Yu, R., Yu, W., &Wang, X. (2024, September). Attention prompting on image for large vision-language models. In*European Conference on Computer Vision* (pp. 251-268). Cham: Springer Nature Switzerland. https://doi.org/10.1007/978-3-031-73404-5_15
- [60] Zhang, J., Zhang, Y., Leach, K., &Huang, Y. (2026). EyeLayer: Integrating Human Attention Patterns into LLM-Based Code Summarization. *arXiv preprint arXiv:2602.22368*. <https://doi.org/10.48550/arXiv.2602.22368>
- [61] Zhao, Y., Wu, H., &Xu, B.(2025, April). Leveraging attention to effectively compress prompts for long-context llms. In *Proceedings of the AAAI Conference on Artificial Intelligence* (Vol. 39, No. 24, pp. 26048-26056). <https://doi.org/10.1609/aaai.v39i24.34800>
- [62] Zheng, Z., Wang, Y., Huang, Y., Song, S., Yang, M., Tang, B., ... &Li, Z. (2025). Attention heads of large language models. *Patterns*, 6(2). DOI: [10.1016/j.patter.2025.101176](https://doi.org/10.1016/j.patter.2025.101176)