

The Integration of Large-Scale Language Models Into Intelligent Adjudication: Justification Rules and Implementation Pathways

SUN, Wenjian ^{1*} WAN, Weixiang ² PAN, Linying ³ XU, Jingyu ⁴ ZENG, Qiang ⁵

¹ Yantai University, China

² University of Electronic Science and Technology of China, China

³ Trine University, USA

⁴ Northern Arizona University, USA

⁵ Zhejiang University, China

* SUN, Wenjian is the corresponding author, E-mail: swjhuman@gmail.com

Abstract: At the end of November 2022, after OpenAI released a new generation of generative artificial intelligence (AIGC) chat tool ChatGPT, it triggered a global network carnival, and articles came to it in terms of questions, conversations, interviews, and reviews. In addition to the noise and uproar, scholars from all walks of life also carried out rational reflection, and the topics of discussion mainly focused on the impact and impact on education, economy, news media, academic research and intellectual property rights and other fields, as well as the exploration of underlying logic issues such as key technologies, operation modes and application scenarios. In essence, large language model (LLMs) is the core framework of a new type of artificial intelligence such as ChatGPT, but the strength of its capability is not only related to corpus, algorithm and computing power, but also closely related to national sovereignty and many security issues, and the biggest obstacle in the realization of artificial intelligence trial is how to realize the legal argument. Therefore, it is necessary to build a set of intelligent trial argumentation rules, and build an algorithm model based on a large language model, to independently complete the trial of the case and output the judgment results for judges' reference, and explore its realization path and apply it to trial practice.

Keywords: Large Language Model, Intelligent Trial, Intelligent Demonstration, Safety Detection

DOI: <https://doi.org/10.5281/zenodo.10607564>

1 Introduction

In recent years, many national courts have been undergoing a new era of digital reform. In the process of the construction of the existing smart court, intelligent trial supervision is still in the initial stage, and the application of information technology still has many pain points to be deepened. The "Research on Key Technologies of Intelligent Trial Supervision in the Whole Process" project aims at the needs of people's courts for big data risk early warning, trial supervision and evaluation in the process of handling cases. This paper studies key technologies such as supervision and early warning of key cases based on multi-dimensional evaluation, identification of major violation handling incidents based on research and judgment, supervision and prediction of important handling nodes based on event chain, and intelligent early warning of key personnel clean government risks based on clues. The project should break through the intelligent computing support method of multi-source data based on supercomputers, develop the supervision and prediction system of important case nodes and the intelligent early

warning system of the risk of integrity of important personnel, build the intelligent trial supervision platform of the whole process, and provide scientific and technological support for the intelligent, integrated and collaborative trial supervision of the court.

However, due to the insufficient, non-objective and distorted case data base, the "black box" effect of the algorithm, the technical bottleneck of artificial intelligence and the shortage of talents, the actual effect of intelligent operation of judicial power is affected and restricted. In the future, the intelligent operation of judicial power in China needs to improve the case database, enhance the explainability of algorithms, strengthen the research and development of intelligent technology, and enhance the application ability and level of intelligent judicial power, so as to truly realize the effective integration of artificial intelligence and law. Based on the discussion and principle analysis of the big language model, this paper analyzes the application path of the big language model in the intelligent trial process from the perspective of intelligent trial.

2 Related Work

As for whether intelligent trial can predict the judgment and the accuracy of predicting the judgment result, some scholars believe that artificial intelligence can extract the case plot through natural semantic recognition technology in criminal cases and deduce the sentencing result according to the previously formed algorithm. The predictions may not be accurate, but as the case moves forward, more information becomes available and the system's sentencing predictions become more accurate.

2.1 Intelligent Trial

With the steady advancement of the information construction of the people's courts, the people enjoy more convenient litigation services, the trial execution assisted by intelligent applications is more efficient, the judicial management activities based on massive data and intelligent analysis are more accurate, and the judicial big data has comprehensively improved the capacity and level of social governance, and the litigation processes such as filing, trial, service and execution are accumulating everywhere. Extremely developed online filing system, electronic service system, voice recognition system, online litigation platform, intelligent assistant case handling system, network investigation and control system, etc., effectively solve the problem of filing difficulty, service difficulty, implementation difficulty and so on. At the same time, a procedural rule system adapted to the Internet era to ensure the orderly online operation of the judiciary has been established. China's court information construction has taken a leading position in the world.

The use of modern scientific and technological means such as big data, cloud computing, and artificial intelligence to solve reform problems, improve judicial efficiency, promote the judicial reform of the people's courts and the construction of intelligence and information technology, and provide strong scientific and technological support for promoting the modernization of the trial system and trial capacity, indicating the direction for the operation of technology-driven trial power. In the future period, the application of intelligent judicial power is reflected in the application of intelligent judicial power system, that is, after the corresponding elements and elements of the case are input into the system, the system can produce certain opinions or suggestions according to the massive data of the case information database and the legal information database for judges to make decisions. In particular, the accurate prediction of the judgment results of conventional cases by artificial intelligence has passed the rule of law awareness and rule of law rules to the public, and played a positive role in good law and good governance. At the same time, it is also necessary to master the use of online simultaneous case handling, trial video, law enforcement records, electronic files and other audio and video materials, the integration of relevant information, the disclosure of trial and service in the application of science and technology

courts, so as to improve the intelligent management information system of judicial power, and free judges from mechanical and repetitive trial activities. Give full play to its active role in interpreting the law and filling the gaps in the law.

2.2 Large Language Model (LLM)

As a proven and feasible direction, the "big" of large language model is reflected in the wide training data set, the large model parameters and layers, and the large calculation amount. Its value is reflected in the universality, and it has better generalization ability. Compared with the traditional language model trained in specific fields, it has a wider range of application scenarios.

Before Transformer was proposed, the mainstream model in the field of natural language processing was recurrent neural network (RNN), which uses recursion and convolutional neural networks for language sequence transformation. In 2017, the Google Brain team's top meeting in the field of artificial intelligence, NeurIPS, published a paper titled "Attention is all you need," which first proposed a new simple network architecture, Transformer, which is based entirely on the attention mechanism. Circular recursion and convolution are completely eliminated.

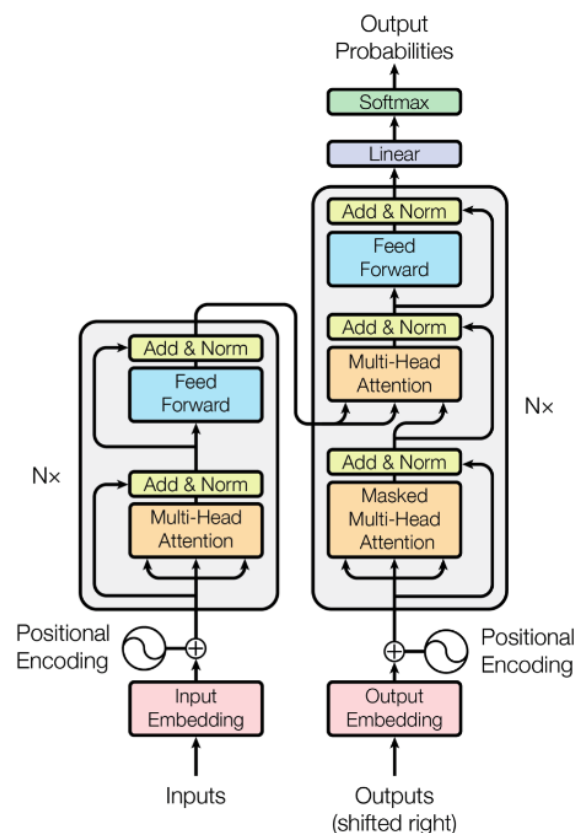


Figure 1: The Transformer - model architecture.

Recursive models typically calculate along the sign positions of input and output sequences to predict

subsequent values. However, this inherent sequential nature prevents parallelization within training samples because memory constraints limit batch processing between samples. The attention mechanism allows dependencies to be modeled without regard to their distance in the input or output sequence.

Transformer eschews the model architecture of a recursive network and relies entirely on attention mechanisms to draw global dependencies between inputs and outputs. After just 12 hours of training on eight P100 Gpus, the Transformer reaches a new state of the art level in terms of translation quality, demonstrating excellent parallelism. It became the most advanced Large Language Model (LLM) at the time.

It breaks through the learning limitations of long-distance text dependence, avoids the model architecture of recursive networks, and completely relies on attention mechanisms to draw global dependencies between inputs and outputs. The operand required to correlate signals from two arbitrary input or output positions, which used to require linear or logarithmic growth as the distance increases, is now converged to a constant, and accuracy is guaranteed by the multi-attention head mechanism. It can also be trained in a high degree of parallel, which is very important for exploiting hardware dividends and quickly iterating models.

The quest for precise robot positioning within logistics automation has been a focal point in recent research endeavors. A multitude of techniques and methodologies have been explored to address this critical challenge.

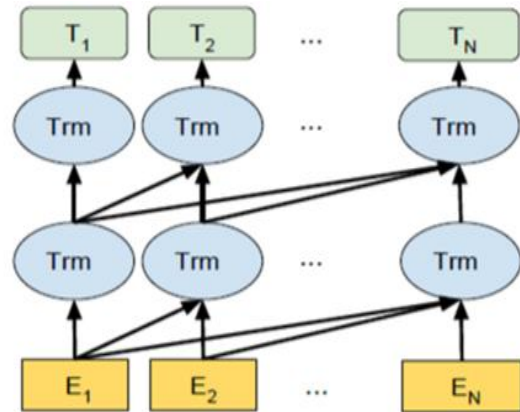


Figure 2: Large language model pre-training

Large language model pre-training is to predict the next word by the words above, which is an unsupervised pre-training. For example, given an unsupervised corpus $U=\{u_1, \dots, u_n\}$, and the pre-trained language model is to maximize the following:

$$L1(U) = \sum_i P(u_i | u_i - k, \dots, u_i - 1; \Theta) \quad (1)$$

As shown in the figure above, to predict the next word through the above, belongs to the autoregressive model, also known as the AR model.

2.3 GPT-3 Model

The GPT-3 model behind the original version of ChatGPT consists of dozens of neural network layers. Each layer takes a series of vectors as input - one for each word in the input text - and adds information to help clarify the meaning of that word and better predict what might come next.

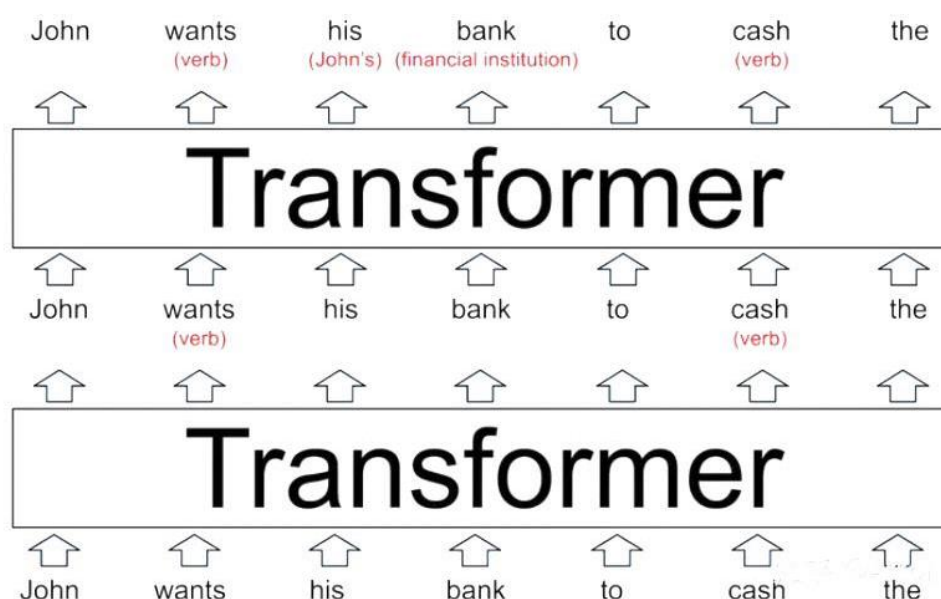


Figure 3: GPT-3 model

It can be seen from the figure that each layer of LLM is a Transformer, and the input text of the model is "John wants his bank to cash the (John wants his bank to cash)", these words are represented as word2vec style vectors. And transfer to the first Transformer. The Transformer determines that both wants and cash are verbs (these words can also be nouns). We represent this additional context with red text in parentheses, but in reality the model stores this information by modifying the word vector in a way that is difficult for humans to interpret. These new vectors are called hidden states and are passed to the next Transformer.

The second Transformer adds two more contextual pieces of information: it clarifies that bank refers to a financial institution, not a riverbank, and that his refers to the pronoun John. The second Transformer generates another set of hidden state vectors that reflect all the information previously learned by the model.

The chart above depicts a purely hypothetical LLM, so don't get too hung up on the details. Real LLMS tend to have more layers. For example, the most powerful version of GPT-3 has 96 layers.

Research shows that the first few layers focus on understanding the syntax of the sentence and resolving the ambiguities shown above. Later layers (the ICONS above are not shown to keep the chart size manageable) are dedicated to a high-level understanding of the entire paragraph. For example, when the LLM "reads" a short story, it seems to remember all sorts of information about the story's characters: gender and age, relationship to other characters, past and current locations, personality and goals, and so on.

The researchers don't fully understand how LLMS keep track of this information, but logically, the model must do so by modifying hidden state vectors when passing time information between layers. The vector dimension in modern LLM is very large, which is conducive to expressing richer semantic information.

Therefore, in combination with the lexical input and output of LLM and Transformer model, this paper analyzes the application scenarios and algorithm processes of large language models in the field of intelligent trial, so as to enrich the intelligent realization of modern trial and argument process with artificial intelligence as the core.

3 Methodology

3.1 Data Processing and Modeling

First, collect legal cases and related legal text data, including case details, legal documents, etc. Data is cleaned, de-duplicated and labeled to ensure data quality and consistency.

Step 2: Using Pre-trained large language models such as BERT (Bidirectional Encoder Representations from Transformers) or GPT (Generative pre-trained Transformer), As a base model.

The data sources required for intelligent trial training of large language model can be divided into general data and professional data. GeneralData includes web pages, books, news, conversation text, and so on. General data has the characteristics of large scale, diversity and easy access, so it can support the construction of language modeling and generalization of large language models.

Specialized Data includes data such as multilingual data, scientific data, code, and domain-specific data. By introducing professional data in the pre-training stage, the task solving ability of large language models can be effectively provided. When modeling language data, as Scatter is very similar to Broadcast, they both have a one-to-many communication mode. The difference is that node 0 of Broadcast sends the same information to all nodes, while Scatter sends different parts of data. Send it to all nodes on demand.

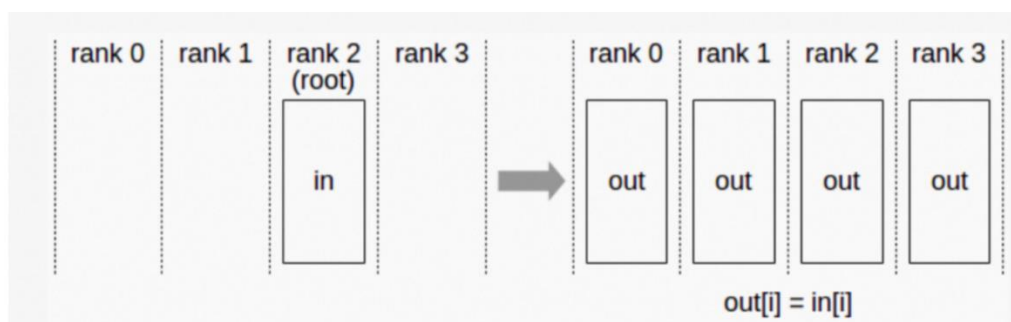


Figure 4: Each information sending node

The sources of scientific text data mainly include arXiv papers, PubMed papers, teaching materials, courseware and teaching web pages. Because the scientific field involves many specialized fields and the data forms are complex, it is usually necessary to preprocess formulas,

chemical formulas, protein sequences, etc., with specific symbolic marks. For example, formulas can be represented using LaTeX syntax, chemical structures can be represented using SMILES (Simplified Molecular Input Line Entry System), and protein sequences can be represented using

single-letter codes or three-letter codes. In this way, data in different formats can be converted into a unified form, making language models better able to process and analyze scientific text data.

Step 3: Customize Transformer structure to adapt to intelligent trial tasks, including input coding, self-attention mechanism, multi-head attention and other components.

3.2 Argument and Judgment

The quality of data on the Internet varies widely, both from OpenAI co-founder Andrej Karpathy's report at Microsoft Build 2023 to some current research showing that the quality of training data has a very important impact on the effectiveness of large language models. Therefore, how to remove low-quality data from the collected data becomes an important step in the training of large language models. The low quality data filtering methods used in large language model training can be roughly divided into two categories: classifier-based methods and heuristic based methods. The goal of the classifier-based approach is to train a text quality judgment model and use the model to identify and filter low-quality data.

GPT-3, PALM and GLam models all use classifier-based methods in training data construction. The Feature Hash Based Linear Classifier can be used to judge the text quality very efficiently. The classifier is trained using a select set of texts (Wikipedia, books, and a few selected websites) with the goal of assigning high scores to web pages that are similar to the training data. This classifier can be used to evaluate the content quality of web pages. In practical applications, it is also possible to select a suitable data set by sampling a web page using Pareto distribution and selecting an appropriate threshold based on its score.

However, some studies have also found that classification-based approaches may remove high-quality text containing dialects or spoken languages, thereby losing some diversity. A heuristic based approach eliminates low-quality text through a carefully designed set of rules, and BLOOM and Gopher take a heuristic based approach. These heuristic rules mainly include:

Language filtering: If a large language model focuses on only one or a few languages, it can significantly filter out text from other languages in the data.

Metrics filtering: Low quality text can also be filtered using metrics. For example, a language model can be used to calculate the Perplexity value for a given text, which can be used to filter out unnatural sentences.

Statistical feature filtering: For text content, statistical features including punctuation distribution, Symbol-to-Word Ratio, sentence length, etc., can be calculated, and low-quality data can be filtered using these features.

Keyword filtering: Based on a specific set of keywords, you can identify and remove noisy or useless elements in the text, such as HTML tags, hyperlinks, and

offensive words.

In the process of data screening for judicial documents of intelligent trial by large language model, the following three steps will be realized.

Step 1: Input legal cases and related information, and let the model automatically argue and decide.

Step 2: The model provides information about the legal basis, jurisprudence and recommendations of the case by analyzing and comparing the case.

Step 3: Use the self-attention mechanism and multi-head attention to capture key information and logical relationships in the text to support the decision process.

The Convolutional Neural Network (CNN) model is a powerful tool for extracting hierarchical features from structured.

3.3 Interpretation and Optimization of Results

In the training and result generation of large-scale language models, there may be hidden algorithm bias or even algorithm discrimination in the whole process of data use, algorithm design, algorithm purpose and algorithm standard determination. In addition, the application and opacity of algorithm rights in the trial will inevitably bring the trust of judges, parties and society. This makes it highly likely that the output of AI trials will not be trusted and used as a reference. Algorithms operate on the basis of case data, which is inherently biased. Therefore, in the process of algorithm operation, these biases will be transmitted to the output result of the algorithm, resulting in the output result is biased. This bias stems from the varying quality of case data, some of which are high quality and some of which are low quality. Intelligent auxiliary machines cannot accurately judge the quality of data, so they will equally learn all case data in the learning process and use it as an important basis for judgment conclusions.

As can be seen from the above method demonstration, this algorithm process combines a large language model with Transformer structure, enabling it to intelligently process legal cases and provide arguments and decision recommendations. The process makes full use of Transformer's self-attention mechanism to deal with complex relationships in text data, which improves the efficiency of the model in intelligent trial.

4 Conclusion

This study explores the key algorithm flow of large language model (LLMs) combined with Transformer in the field of intelligent trial, and analyzes its application potential in intelligent legal case argument and decision. Using a combination of pre-trained LLMs and a custom Transformer structure, we are able to automate legal case analysis, argument and decision generation, providing judges with strong decision support. With further improvements in data

quality and algorithmic transparency, this approach is expected to play a greater role in the field of intelligent trials.

The algorithm-based process combined with Large Language models (LLMs) and Transformer offers new prospects for intelligent trial and case argument. In the future, we can expect more courts to adopt this technology to improve the efficiency and accuracy of decisions. At the same time, with improved data quality and reduced algorithmic bias, LLMs and Transformer will become powerful tools in the legal field to support fair, impartial and efficient justice. The field remains challenging, but we are confident in its future and believe that this technology will continue to evolve and refine, bringing more innovation and progress to smart trials and legal arguments.

Acknowledgments

Finally, I would like to extend my heartfelt thanks to Professor Yiming Pan, whose research has provided valuable reference and inspiration in the writing of this article. In particular, his article "The implementation of an AI-driven advertising push system based on a NLP algorithm" provides important support and guidance for the relevant content of this paper. I would like to express my deep gratitude for his remarkable work in the fields of natural language processing and artificial intelligence, as well as for his contributions to this article.

Funding

Not applicable.

Institutional Review Board Statement

Not applicable.

Informed Consent Statement

Not applicable.

Data Availability Statement

The original contributions presented in the study are included in the article/supplementary material, further inquiries can be directed to the corresponding author.

Conflict of Interest

The authors declare that the research was conducted in the absence of any commercial or financial relationships that could be construed as a potential conflict of interest.

Publisher's Note

All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers. Any product that may be evaluated in this article, or claim that may be made by its manufacturer, is not guaranteed or endorsed by the publisher.

Author Contributions

Not applicable.

About the Authors

SUN, Wenjian

Postgraduate, Electronic and information engineering, Yantai University, ShanDong.

WAN, Weixiang

Postgraduate, Electronics & Communication Engineering, University of Electronic Science and Technology of China, Chengdu, China.

PAN, Linying

Postgraduate, Information studies, Trine University, Phoenix, Arizona, USA.

XU, Jingyu

Postgraduate, Computer Information Technology, Northern Arizona University, Flagstaff, Arizona, USA.

ZENG, Qiang

Postgraduate, Computer Technology, Zhejiang University, Hangzhou, Zhejiang, China.

References

- [1] LIU X, HE P, CHEN W, et al. Multitask deep neural networks for natural language understanding[C]//Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics, Florence, Italy:2019:Association for Computational Linguistics,4487-4496.
- [2] MIKOLOV T, CHEN K, CORRADO G, et al. Efficient estimation of word representations in vector space[C]//Proceedings of the International Conference on Learning Representations (ICLR 2013), Scottsdale,AZ, 2013.1-12.DOI: 10.48550/arXiv.1301.3781.
- [3] Liu, B. (2023). Based on intelligent advertising recommendation and abnormal advertising monitoring system in the field of machine learning. International Journal of Computer Science and Information Technology, 1(1), 17-23.

- [4] Pan, Linying, et al. "Research Progress of Diabetic Disease Prediction Model in Deep Learning". *Journal of Theory and Practice of Engineering Science*, vol. 3, no. 12, Dec. 2023, pp. 15-21, doi:10.53469/jtpes.2023.03(12).03.
- [5] K. Tan and W. Li, "Imaging and Parameter Estimating for Fast Moving Targets in Airborne SAR," in *IEEE Transactions on Computational Imaging*, vol. 3, no. 1, pp. 126-140, March 2017, doi: 10.1109/TCI.2016.2634421.
- [6] Zhou, H., Lou, Y., Xiong, J., Wang, Y., & Liu, Y. (2023). Improvement of Deep Learning Model for Gastrointestinal Tract Segmentation Surgery. *Frontiers in Computing and Intelligent Systems*, 6(1), 103-106.
- [7] Liu, S., Wu, K., Jiang, C., Huang, B., & Ma, D. (2023). Financial Time-Series Forecasting: Towards Synergizing Performance And Interpretability Within a Hybrid Machine Learning Approach. *arXiv preprint arXiv:2401.00534*.
- [8] Su, J., Nair, S., & Popokh, L. (2023, February). EdgeGym: A Reinforcement Learning Environment for Constraint-Aware NFV Resource Allocation. 2023 IEEE 2nd International Conference on AI in Cybersecurity (ICAIC), 1–7. doi:10.1109/ICAIC57335.2023.10044182
- [9] Su, J., Nair, S., & Popokh, L. (2022, November). Optimal Resource Allocation in SDN/NFV-Enabled Networks via Deep Reinforcement Learning. 2022 IEEE Ninth International Conference on Communications and Networking (ComNet), 1–7. doi:10.1109/ComNet55492.2022.9998475
- [10] Popokh, L., Su, J., Nair, S., & Olinick, E. (2021, September). IllumiCore: Optimization Modeling and Implementation for Efficient VNF Placement. 2021 International Conference on Software, Telecommunications and Computer Networks (SoftCOM), 1–7. doi:10.23919/SoftCOM52868.2021.9559076
- [11] Bao, W., Che, H., & Zhang, J. (2020, December). Will_Go at SemEval-2020 Task 3: An Accurate Model for Predicting the (Graded) Effect of Context in Word Similarity Based on BERT. In A. Herbelot, X. Zhu, A. Palmer, N. Schneider, J. May, & E. Shutova (Eds.), *Proceedings of the Fourteenth Workshop on Semantic Evaluation* (pp. 301–306). doi:10.18653/v1/2020.semeval-1.
- [12] WANG N, LI S L, LIU T L, et al. Attention basedBiGRU on tendency analysis of judgment results[J].*Computer Systems & Applications*, 2019, 28 (3):191-195.DOI: 10.15888/j.cnki.csa.006816.
- [13] Xinyu Zhao, et al. "Effective Combination of 3D-DenseNet's Artificial Intelligence Technology and Gallbladder Cancer Diagnosis Model". *Frontiers in Computing and Intelligent Systems*, vol. 6, no. 3, Jan. 2024, pp. 81-84, <https://doi.org/10.54097/iMKyFavE>.
- [14] Shulin Li, et al. "Application Analysis of AI Technology Combined With Spiral CT Scanning in Early Lung Cancer Screening". *Frontiers in Computing and Intelligent Systems*, vol. 6, no. 3, Jan. 2024, pp. 52-55, <https://doi.org/10.54097/LAwfJzEA>.
- [15] Liu, Bo & Zhao, Xinyu & Hu, Hao & Lin, Qunwei & Huang, Jiaxin. (2023). Detection of Esophageal Cancer Lesions Based on CBAM Faster R-CNN. *Journal of Theory and Practice of Engineering Science*. 3. 36-42. 10.53469/jtpes.2023.03(12).06.
- [16] Yu, Liqiang, et al. "Research on Machine Learning With Algorithms and Development". *Journal of Theory and Practice of Engineering Science*, vol. 3, no. 12, Dec. 2023, pp. 7-14, doi:10.53469/jtpes.2023.03(12).02.
- [17] Xin, Q., He, Y., Pan, Y., Wang, Y., & Du, S. (2023). The implementation of an AI-driven advertising push system based on a NLP algorithm. *International Journal of Computer Science and Information Technology*, 1(1), 30-37.0
- [18] Zhou, H., Lou, Y., Xiong, J., Wang, Y., & Liu, Y. (2023). Improvement of Deep Learning Model for Gastrointestinal Tract Segmentation Surgery. *Frontiers in Computing and Intelligent Systems*, 6(1), 103-106.6
- [19] Implementation of an AI-based MRD Evaluation and Prediction Model for Multiple Myeloma. (2024). *Frontiers in Computing and Intelligent Systems*, 6(3), 127-131. <https://doi.org/10.54097/zJ4MnbWW>.
- [20] Zhang, Q., Cai, G., Cai, M., Qian, J., & Song, T. (2023). Deep Learning Model Aids Breast Cancer Detection. *Frontiers in Computing and Intelligent Systems*, 6(1), 99-102.3
- [21] Xu, J., Pan, L., Zeng, Q., Sun, W., & Wan, W. (2023). Based on TPUGRAPHS Predicting Model Runtimes Using Graph Neural Networks. *Frontiers in Computing and Intelligent Systems*, 6(1), 66-69.7
- [22] Wan, Weixiang, et al. "Development and Evaluation of Intelligent Medical Decision Support Systems." *Academic Journal of Science and Technology* 8.2 (2023): 22-25.
- [23] Tian, M., Shen, Z., Wu, X., Wei, K., & Liu, Y. (2023). The Application of Artificial Intelligence in Medical Diagnostics: A New Frontier. *Academic Journal of Science and Technology*, 8(2), 57-61.7
- [24] Shen, Z., Wei, K., Zang, H., Li, L., & Wang, G. (2023). The Application of Artificial Intelligence to The Bayesian Model Algorithm for Combining Genome Data. *Academic Journal of Science and Technology*, 8(3), 132-135.2
- [25] VASWANI A, SHAZEER N, PARMAR N, et al. Attention is all you need [C]//*Advances in Neural Information Processing Systems*. Long Beach,

CA.USA: NIPS, 2017: 5998-6008

- [26] Du, Shuqian, et al. "Application of HPV-16 in Liquid-Based Thin Layer Cytology of Host Genetic Lesions Based on AI Diagnostic Technology Presentation of Liquid". *Journal of Theory and Practice of Engineering Science*, vol. 3, no. 12, Dec. 2023, pp. 1-6, doi:10.53469/jtpes.2023.03(12).01.
- [27] Pan, Yiming, et al. "Application of Three-Dimensional Coding Network in Screening and Diagnosis of Cervical Precancerous Lesions". *Frontiers in Computing and Intelligent Systems*, vol. 6, no. 3, Jan. 2024, pp. 61-64, <https://doi.org/10.54097/mi3VM0yB>.
- [28] Song, Tianbo, et al. "Development of Machine Learning and Artificial Intelligence in Toxic Pathology". *Frontiers in Computing and Intelligent Systems*, vol. 6, no. 3, Jan. 2024, pp. 137-41, <https://doi.org/10.54097/Be1ExjZa>.
- [29] Qian, Jili, et al. "Analysis and Diagnosis of Hemolytic Specimens by AU5800 Biochemical Analyzer Combined With AI Technology". *Frontiers in Computing and Intelligent Systems*, vol. 6, no. 3, Jan. 2024, pp. 100-3, <https://doi.org/10.54097/qoseeQ5N>.
- [30] Zheng, Jiajian, et al. "The Credit Card Anti-Fraud Detection Model in the Context of Dynamic Integration Selection Algorithm". *Frontiers in Computing and Intelligent Systems*, vol. 6, no. 3, Jan. 2024, pp. 119-22, <https://doi.org/10.54097/a5jafgdv>.