

# ***Research Analysis on Adaptive Thought Chains Based on Knowledge Distillation***

**Yushan Xia**

*School of Information Engineering, Nantong Institute of Technology, Nantong, China  
2310410081@ntit.edu.cn*

**Abstract.** Large-scale language models (LLMs) have a huge scale, but they have encountered difficulties in improving social productivity, mainly due to their high cost and increasing demand for a large number of computing resources. Knowledge refining plays a key role in bridging the gap between model performance and operational efficiency, and also strengthens these two aspects. This technology refines the functions of GPT-3.5 and other models into miniature models that can be run locally at controllable cost. It can not only enable small and medium-sized enterprises and research institutions to use high-performance large-scale language models, but also ensure data security. Classical knowledge extraction framework uses label softening to achieve knowledge transfer between teacher and student models. According to the logical steps of teacher-student model alignment (similar to GPT-4), this paper mainly focuses on making the student model learn from the teacher model adaptively. This method enables the student model to get a compact model quickly, which not only absorbs the fine knowledge of the teacher model, but also reduces the consumption of computing resources and data. The existing CoT distillation methods ignore the variability of samples and the learning dynamics of student models. In this paper, an adaptive chain distillation method is proposed, so that small models can avoid the problem of reasoning sensitivity and focus on learning difficult samples. Although this will weaken its ability to analyze complex problems, we introduce an adaptive reasoning mechanism including soft prompt fine-tuning module and do experiments to verify it.

**Keywords:** Knowledge distillation, Large language models, Teacher-student model, Chain of reasoning

## **1. Introduction**

Large language model technology has advanced at an astonishing pace in recent years, with models like GPT-3.5 now boasting parameter scales reaching the tens of billions. Though their training and inference processes consume substantial computational resources, they demonstrate exceptional reasoning capabilities across complex tasks in multiple domains. Achieving this requires models to possess vast knowledge reserves and employ multi-level reasoning to devise solutions for uncertain problems. While today's large language models possess computational power and advanced reasoning capabilities that surpass human performance, they face significant resource constraints—exorbitant computational costs make them unaffordable for small and medium-sized

enterprises. During this period, market demand for models with complex reasoning capabilities has grown steadily—capabilities required in fields such as industrial fault diagnosis and scientific literature analysis. Among them, multi-level knowledge integration and uncertain solution are the most important. In this case, it is very important to give small models like BERT-base and DiTI-LGP2 complex reasoning ability to promote artificial intelligence technology. In order to endow the small model with complex reasoning ability, the "fine chain thinking reasoning" method optimizes the small model with the chain thinking reasoning ability of the large model. Regarding the knowledge distillation technology, Wang and other scholars put forward a method of reasoning knowledge distillation, which achieves the goal of collaborative optimization by coordinating the integration of the reasoning stage of the master model and the slave model [1]. This method enables the small model to perform as well as the huge language model with 7 billion parameters in reasoning tasks. However, the existing technology still has limitations: most methods rely too much on the static data in large models, and it is difficult to adapt to complex and ever-changing tasks; Insufficient reasoning depth will reduce performance.

This paper addresses the enhancement of complex reasoning capabilities in small models by proposing a method to distill this chain of reasoning technology—which incorporates awareness of problem complexity—into compact models. It introduces a perception module that prompts small models to consider the complexity of input problems.

## 2. Overview of the development of thought chain distillation technology

### 2.1. Technical architecture

Knowledge distillation methods further integrate contrastive learning and multimodal approaches to enhance student models.

The vector of BERT-base is used for multimodal feature extraction, the last layer feature of ResNet50 is used for image data, and MFCC feature is combined with transformer encoder for voice data. Different modes are fused into a unified cross-modal feature by means of modal attention weighting. Compared with the traditional text extraction, the multimodal method enables the student model to deal with complex reasoning tasks including text and images.

CoT method mainly improves the reasoning ability of large-scale models in complex tasks, and uses hundreds of millions of parameters. But it can't improve the reasoning ability of small models. In order to make the small model perform complex reasoning, efforts have been made to improve the reasoning data quality of the large model based on the method of fine-tuning CoT [3]. Using the idea of chain, representative questions are used to generate reasoning processes and answers, but the adaptability is still lacking [4].

Adaptive computational models can perceive and evaluate problem complexity while accomplishing problem-solving tasks.

### 2.2. Core technical features

The teacher model utilizes the GPT-3.5 architecture, with its reasoning process relying on API calls to OpenAI. In traditional Fine-tune-CoT approaches, the teacher model generates a reasoning chain by inputting questions via the API, then combines the question and reasoning to produce the answer. In this experiment, during the initial training phase, the teacher model provides detailed foundational reasoning chains to pre-train the smaller model [5]. To improve adaptive capabilities, dynamic reasoning is necessary to enhance reasoning efficiency.

### 3. Common datasets and evaluation criteria

#### 3.1. Core dataset classification

The AddSub and Date Understanding datasets were selected as testing platforms [6]. AddSub is a mathematical dataset primarily used for natural language reasoning tasks, requiring the system to perform arithmetic operations and infer answers. The Date Understanding dataset, on the other hand, focuses on temporal comprehension in natural language processing, assessing and processing date representations in various formats. Through in-depth analysis, the efficiency and accuracy of the teacher model's generation are enhanced, thereby providing high-quality reasoning chains for subsequent training of the student model.

We selected symbol-based tasks (Last Letter Concatenation), arithmetic tasks (SingleEq, MultiArith, AddSub), common sense tasks (CommonSenseQA, StrategyQA), and logic tasks (LogiQA) to compare with static CoT distillation, thereby demonstrating the effectiveness of the proposed method.

#### 3.2. Evaluation system

The teacher model, InstructGPT, was modified following the deprecation of OpenAI's text-davinci-001, text-davinci-002, and other similar models [7]. The experimental datasets used are AddSub and Date Understanding. In contrast to the original Fine-tune-CoT method, which directly generated reasoning data using a zero-shot-CoT approach, the new method generates reasoning data for all questions using zero-shot-CoT, and then fine-tunes the student model. Accuracy on the AddSub dataset showed a notable improvement, increasing from 76.71% to 83.54%. On the Date Understanding dataset, the initial accuracy was also 76.71%. Through heuristic error correction on the teacher model and fine-tuning the student model, accuracy was improved to 78.88% [8]. However, whether the student model is trained by enhancing the teacher model's performance or by adapting to perceptual issues, the loss gap between correct and incorrect samples remains small. As training progresses and the student model engages in reasoning, this gap gradually widens. This suggests that the model places greater emphasis on learning from incorrect samples to improve its overall performance.

### 4. Distilling knowledge at the current state of technology

#### 4.1. Core achievements

The core focus of current technology is "adaptivity." Whether adjusting the length of the inference chain or refining the training method, the goal is to optimize the distillation process to align precisely with the existing capabilities of the student model. High-quality, small-scale datasets built using adaptive technology often enhance the performance of student models more effectively than large, redundant, and disorganized data. Furthermore, distillation technology is expanding into fields such as medicine, mathematics, and multimodal applications, reducing costs and improving computational efficiency, thus demonstrating its broad applicability.

#### 4.2. Direction of development

The development of adaptive chain distillation technology mainly focuses on three aspects. In terms of technical improvement and innovation, more complex adaptive mechanisms will be established,

advanced teaching models integrating new paradigms will be developed, and small models will be endowed with reasoning ability [9]. Expanding the application field will promote the application of this technology to a variety of scenarios, including multimodal reasoning and vertical industrial applications, and make contributions to many fields and aspects. The technology application field attaches importance to the use of acceleration technology to optimize process automation, while using ultra-light hardware across technical fields. Through the coordinated development of these three fields, adaptive chain distillation technology will be fully developed and applied in practice.

## 5. Issues with mainstream methods and their solutions

### 5.1. Core issue

In the knowledge distillation of CoT, the existing methods rely too much on the fixed paradigm, which has caused some obvious limitations [10]. These methods ignore the difference of sample complexity. In simple problems, redundant reasoning chains will waste resources. In complex problems, an oversimplified chain can't give enough guidance. This method lacks flexibility and can't meet the requirements of dynamic changes of student models. In the early training stage, the too complicated reasoning chain is beyond the processing ability of the student model, and in the later stage, the too simplified chain will reduce the learning efficiency. In addition, the existing methods can't distinguish the quality of the teacher model reasoning chain, which leads to the noise pollution of the student model.

In order to meet these challenges, this paper proposes an adaptive chain distillation method. This method adopts a dynamic adaptive mechanism: a complexity module according to the difficulty of the task will dynamically change the detail and length of the reasoning chain. At the beginning of the model, the reasoning steps will be complicated, but with the deepening of training, the model will develop towards simplification, thus reducing the interference caused by redundant information. This method attempts to obtain a compact model with powerful performance, efficient reasoning speed and excellent generalization ability by distillation, so as to overcome the limitations of mainstream methods.

### 5.2. Tailored solutions

In order to overcome the above problems and improve the accuracy and reliability of the results, this paper combines selective intervention with highly reliable methods to dynamically correct the output path. Real-time monitoring mechanism corrects the deviation immediately after finding the deviated nodes in the process of model generation to prevent the deviation from accumulating and spreading. Using the highly reliable main model output as a benchmark, the algorithm can filter out long tail noise data such as low probability misclassification and non-critical scene interference, ensure the accuracy of the results, and minimize the negative impact of invalid output on practical application efficiency [10].

In order to reduce the cost and improve the compatibility, this topic only optimizes the compression of redundant data in the reasoning chain, reducing repeated calculation and eliminating redundant intermediate parameters. This method can greatly reduce the overall calculation load during the experiment and accelerate the response speed of the model. In order to deal with the compatibility problem, we adopt the optimal migration theory and align the algorithms in vocabulary distribution, thus solving the problems of vocabulary system differences and data format

conflicts between different technical architectures. These improvements eliminate architectural obstacles and realize seamless interoperability and fusion of data across multiple systems.

This experiment uses a learning model with adaptive difficulty, which can gradually acquire complex abilities. Build a basic model with simple tasks and make initial improvement. After the basic task is completed, the complexity gradually increases, from single scene integration to multi-scene integration, from basic operation to deep logical reasoning. The core competence of teaching mode can be accurately transmitted across different levels, which ensures multi-scenario adaptation when dealing with complex tasks. Progressive training makes the model understand the underlying logic of complex tasks slowly, and makes the model balance comprehensiveness and practicality in the process of ability transfer.

## 6. Conclusion

The adaptive reasoning chain mechanism proposed in this paper condenses the reasoning ability of large model into compact model, which makes the latter form adaptive thinking ability through training. This greatly improves the ability of compact model to solve complex problems autonomously. A detailed comparison across data sets shows that the adaptive inference chain is more feasible than the traditional distillation method. The addition of perception module strengthens the adaptive dynamic reasoning ability of small model. The innovative mechanism improves the performance of the small model, and has obvious advantages in optimization, which is suitable for the environment with limited resources. The integration of automation modules will become an important force to promote this technology, so that the mechanism of adaptive reasoning chain can be better implemented, thus meeting the increasingly complicated task requirements.

In practice, the adaptive chain reasoning mechanism has strong universality and expansibility, and can flexibly deal with reasoning tasks in different fields and scenarios. Its advantages will support more complicated reasoning and decision-making process, so as to fit in with the constantly developing technology and adapt to the vast application situation. Combining this mechanism with other cutting-edge technologies may become an important research hotspot. Integrated reinforcement learning technology can further improve the adaptability and decision-making ability of compact model in complex environment, and make the mechanism more robust and flexible in dynamic scene.

In order to improve the output accuracy, the future research will focus on the control and optimization of each key component in the reasoning chain. Design accurate algorithms and strategies to ensure the reliability and stability of the final output results. By carefully adjusting each reasoning step and modeling accurately, the performance of small model can be further improved when dealing with complex reasoning tasks.

As far as distillation efficiency is concerned, more effective distillation algorithm and architecture design are studied to reduce the consumption of computing resources and improve the processing speed, especially for a large number of data and complex operation tasks. Using optimized distillation workflow and hardware collaborative design, researchers should consider the responsiveness and processing efficiency of the system on the premise of ensuring the quality of reasoning. The future work will focus on building a powerful transmission framework and evaluation system to measure the performance of compact models in complex tasks. This will help developers better understand various transfer strategies and lay a solid foundation for technical iteration. Through continuous research and innovation, the adaptive reasoning chain mechanism will play an increasingly important role, promote the development of artificial intelligence, and provide more efficient intelligent solutions for various industries.

## References

- [1] Wang, L., Yoon, K. J. (2021). Knowledge distillation and student-teacher learning for visual intelligence: A review and new outlooks. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 44(6): 3048-3068.
- [2] Liu, F. (2025). Research on Chinese Spelling Correction Methods Based on Multimodal Feature Fusion. North China University of Technology. DOI: 10.26926/d.cnki.gbfgu.2025.000686.
- [3] Brown, T., Mann, B., Ryder, N. et al., (2020). Language mod-els are few-shot learners, " *Advances in neural information processing systems*, vol. 33, pp. 1877-1901.
- [4] Zhang, Z., Zhang, A., Li, M., et al. (2022). Automatic chain of thought prompting in large language models. <https://arxiv.org/abs/2210.03493>.
- [5] Huang, J., Gu, S. S., Hou, L., et al. (2022). Large language models can self-improve. <https://arxiv.org/abs/2210.11610>.
- [6] Szegedy, C., Vanhoucke, V., Ioffe, S., et al. (2016). Rethinking the inception architecture for computer vision. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2818-2826.
- [7] Zhang, Y., Xiang, T., Hospedales, T. M., et al. (2018). Deep mutual learning. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 4320-4328.
- [8] Chu, Z., Chen, J., Chen, Q. et al., (2023). A survey of chain of thought reasoning: Advances, frontiers and future, *arXiv preprint arXiv: 2309.15402*.
- [9] Ji, X. L. (2025). Dynamic Self-Optimization: An Adaptive Standard and Feedback-Driven Optimization Framework for Large Language Model Question-Answering. *Intelligent Computer and Applications*, 1-8. <https://doi.org/10.20169/j.issn.2095-2163.25072303>.
- [10] Ding, Y., Chang, J., Liu, Y. M., et al. (2025). Knowledge Distillation for Efficient Deployment and Application of Large Language Models. *Information and Communication Technology*, 19(03): 53-60. DOI: CNKI: SUN: OXXT.0.2025-03-008.