

Bosch@AI_Team at LegalSML 2025: Vietnamese Legal Small Language with Domain Adaptation and Aspect-based Data Synthesis

Tran Minh Quang^{1,2}, Nguyen Xuan Phi¹, Nguyen Van Tai¹, Phan Minh Toan¹, Dang Van Thin³

¹Bosch Global Software Technologies, Ho Chi Minh City, Vietnam

²Royal Melbourne Institute of Technology Vietnam, Ho Chi Minh City, Vietnam

³University of Information Technology-VNUHCM, Ho Chi Minh City, Vietnam

{quang.tranminh2,tai.nguyenvan,toan.phanminh}@vn.bosch.com

s3757281@rmit.edu.vn

external.Phi.NguyenXuan@bcn.bosch.com

thindv@uit.edu.vn

Abstract

This paper presents our solution designed for LegalSML of VLSP 2025 shared task, which challenged participants to develop small-to-medium sized language models for the Vietnamese legal domain. To address this, we developed a framework that combines synthetic instruction tuning with chain-of-thought question-answer data to guide a large teacher model. This data was then used with both Parameter-Efficient Fine-Tuning (PEFT) and Full-parameter Fine-Tuning (FPFT) to adapt the teacher model into a compact student model optimized for the Vietnamese legal domain. Our approach secured the top-1 average score across the three core evaluation tasks and achieved highly competitive results in both the NLI and Multiple Choice QA tasks.

1 Introduction

Large Language Models (LLMs) have recently shown impressive abilities in natural language processing (NLP) and other areas (Naveed et al., 2025). With the development of various LLMs such as ChatGPT, DeepSeek (Liu et al., 2024), Qwen (Yang et al., 2025), etc., there is a growing need to apply them to different domains such as medicine (Liu et al., 2025), law (Marcos, 2025), and finance (Lee et al., 2025). Additionally, research on creating specialized language models, especially small language models with fewer parameters but high efficiency, has been a new trend recently (Zhang et al., 2025). The reason for researching small language models is to reduce the costs of AI applications, such as deployment, execution, and inference costs. However, currently, most research focuses on resource-rich languages like English or Chinese, while resource-scarce languages like Vietnamese have not received much attention (Wang et al., 2025).

To address the research gap in lower-resource

languages, the VLSP 2025 organization¹ has launched a challenge on Vietnamese Legal Small Language Models (LegalSLM). The shared task requires participants to develop small to medium-sized language models specialized for the Vietnamese legal domain. The models from participating teams will be evaluated on three different tasks, including Legal Citation Usefulness, Multiple-Choice Legal QA, and Free-Text Legal QA. The details of the core evaluation tasks can be presented on this page². The objective of this competition is to leverage continual pretraining, fine-tuning, and instruction tuning techniques to develop small language models. All models from participating teams must have a maximum capacity of 4 billion parameters.

In this paper, we present our solution in developing small language models that are specialized in the legal domain in the LegalSLM shared task.

Our contributions can be summarized as follows:

- We address the lack of a systematic methodology for generating high-quality synthetic instruction-following data to fine-tune smaller LLMs in the Vietnamese legal domain³.
- We propose a framework that integrates synthetic instruction tuning with chain-of-thought (CoT) question-answer data, combined with parameter-efficient fine-tuning (PEFT) and Full-parameter Fine-Tuning (FPFT), to adapt a large teacher model into a compact student model specialized for Vietnamese legal domain.
- Experimental results demonstrate that our approach outperforms other participants in the average scores of three evaluation tasks and

¹<https://vlsp.org.vn/vlsp2025/>

²<https://vlsp.org.vn/vlsp2025/eval/legalslm>

³https://huggingface.co/datasets/QuangTran276/new_reasoning

achieved the top-1 ranking in the LegalSLM 2025 shared task.

2 Related work

Vietnamese QA and Instruction-Tuned LLMs

Early progress on Vietnamese QA is exemplified by ViGPTQA (Nguyen et al., 2023), which introduced an instruction-tuned Vietnamese model and the ViTruthfulQA benchmark to evaluate truthfulness. This work highlights the importance of native evaluation resources and domain-specific adaptation. Another work (Nguyen et al., 2025) provides the first comprehensive Vietnamese legal QA dataset, containing over 3,000 professionally annotated questions. It highlights unique challenges for general-purpose LLMs, including complex statutory phrasing and citation styles.

Domain Specialization of LLMs Surveys suggest that domain specialization—through domain-adaptive pre-training, supervised fine-tuning, and PEFT—is key for LLMs to achieve impact in specialized contexts (Ling et al., 2023). These insights guide the adaptation of models for Vietnamese legal QA. Traditional full fine-tuning (FFT) approaches (Ouyang et al., 2022), while capable of adapting models to specific domains, are resource-intensive and prone to catastrophic forgetting, where the model loses its general-purpose capabilities as it specializes. To address these limitations, Parameter-Efficient Fine-Tuning (PEFT) techniques have gained prominence. Methods such as Low-Rank Adaptation (LoRA) (Hu et al., 2022) and its quantized variant QLoRA (Dettmers et al., 2023) substantially reduce computational and memory requirements by freezing the majority of the base model’s weights and training only a small subset of parameters. This approach enables the creation of lightweight, domain-specific adapters for multiple tasks without the need to maintain separate full model copies.

3 Methodology

3.1 Data Collection and Preprocessing

A comprehensive Vietnamese legal corpus will be compiled for the instruction tuning in the Section 3.4. This corpus will consist of:

- **Vietnamese Legal Corpus:** A collection of official legal documents, including codes, statutes, decrees, and circulars. These doc-

uments will be sourced from official government websites and legal databases.

- **Legal News and Articles:** To supplement the formal legal texts, we will incorporate legal news, commentaries, and articles from reputable online legal journals and news outlets. This will expose the model to a broader range of legal discourse and contemporary legal issues.

The collected texts will undergo a series of pre-processing steps, including text cleaning to remove irrelevant artifacts (e.g., HTML tags), normalization of text, and sentence segmentation to create a high-quality training dataset.

3.2 Instruction Tuning Dataset Construction

Given the scarcity of labeled data for specialized legal tasks in Vietnamese, we will generate a synthetic dataset for instruction tuning. The use of powerful LLMs for synthetic data generation is an effective strategy to create diverse and high quality data sets on scale (Zhou et al., 2023). Our data generation process will be aspect-based and leverage Chain-of-Thought (CoT) prompting to create structured and explainable training examples. The overall dataset construction is illustrated in Figure 2

3.2.1 Aspect-based Synthetic Data Generation

We adopt an aspect-based pipeline to ensure that synthetic data captures both the content and reasoning structures of Vietnamese legal texts. Instead of treating a document as a single unit, the process decompose it into key aspects and generates examples per aspect, improving diversity. To ensure the quality, we also employ a separate LLM-as-a-judge mechanism to conduct quality verification. The prompt used for the LLM (GPT-4) to extract aspects from the legal texts is shown in the Appendix B.1.

3.2.2 Chain-of-Thought Prompting for Structured and Explainable Reasoning

Based on the aspects identified in Section 3.2.1, the teacher model (GPT-4) is prompted to produce intermediate reasoning steps. These prompts are provided in Appendix B.2 and include structured instructions for multiple-choice and citation tasks, as well as detailed analyses for free-text questions. This design allows the student model not only to generate final answers but also to approximate the

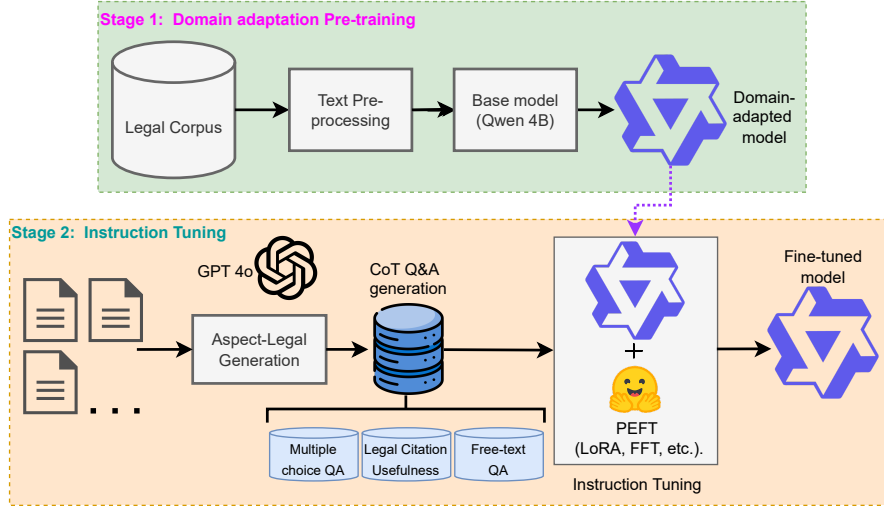


Figure 1: Illustration of our approach to develop specialized Vietnamese legal language model based on the Qwen foundation model.

underlying legal reasoning processes. Using these prompts, we construct Chain-of-Thought (CoT) augmented examples for three representative legal tasks:

- **Legal Citation Usefulness:** A binary classification task where the model determines whether a cited provision addresses a posed legal question, with reasoning explaining the relevance or irrelevance.
- **Multiple-Choice Legal QA:** Aspect-driven multiple-choice questions with four options, including distractors accompanied by reasoning clarifying why the correct answer holds.
- **Free-Text Legal QA:** Open-ended legal questions requiring narrative responses. Answers follow a structured reasoning format with explicit major premise, minor premise, and conclusion.

3.3 Stage 1: Domain adaptation Pre-training

The continual pretraining stage (Figure 1), provided by the shared task organizers, is designed to adapt the base model to the vocabulary, syntax, and semantic characteristics of the Vietnamese legal domain. For this research, an open-source version of the Qwen model (e.g., Qwen-4B) is selected as the foundation, reflecting the emphasis on computationally efficient adaptation. The strong multilingual performance of Qwen models across diverse benchmarks further supports their suitability for subsequent domain-specific fine-tuning.

3.4 Stage 2: Instruction Tuning

The objective of the instruction tuning stage is to adapt the domain-adapted Qwen backbone so that it reliably follows task-level legal instructions and reproduce reasoning behavior exposed in the instruction dataset from the Section 3.2.2. We adopt a parameter-efficient adapter strategy that preserves the frozen pre-trained weights and concentrates optimizing on a small set of newly introduced parameters. The overall pipeline-converting examples into the model chat format, inserting low-rank adapters, optimizing only adapter parameters, and persisting adapter artifacts-is implemented with careful consideration of memory efficiency, computational optimization, and training stability as described in the Section 4.1.

Conversational integration and target formatting:

Instruction examples undergo a comprehensive transformation process to align with the Qwen chat template structure. Each legal instruction-response pair is systematically converted into a conversational format where the model is conditioned on an instruction/context block and trained to generate the appropriate response, including intermediate reasoning steps when available in the source data. During this conversion process, several critical pre-processing steps are implemented: (1) metadata annotations and irrelevant formatting markers are systematically removed to focus the model’s attention on core legal content; (2) multi-turn conversational inputs are normalized to ensure

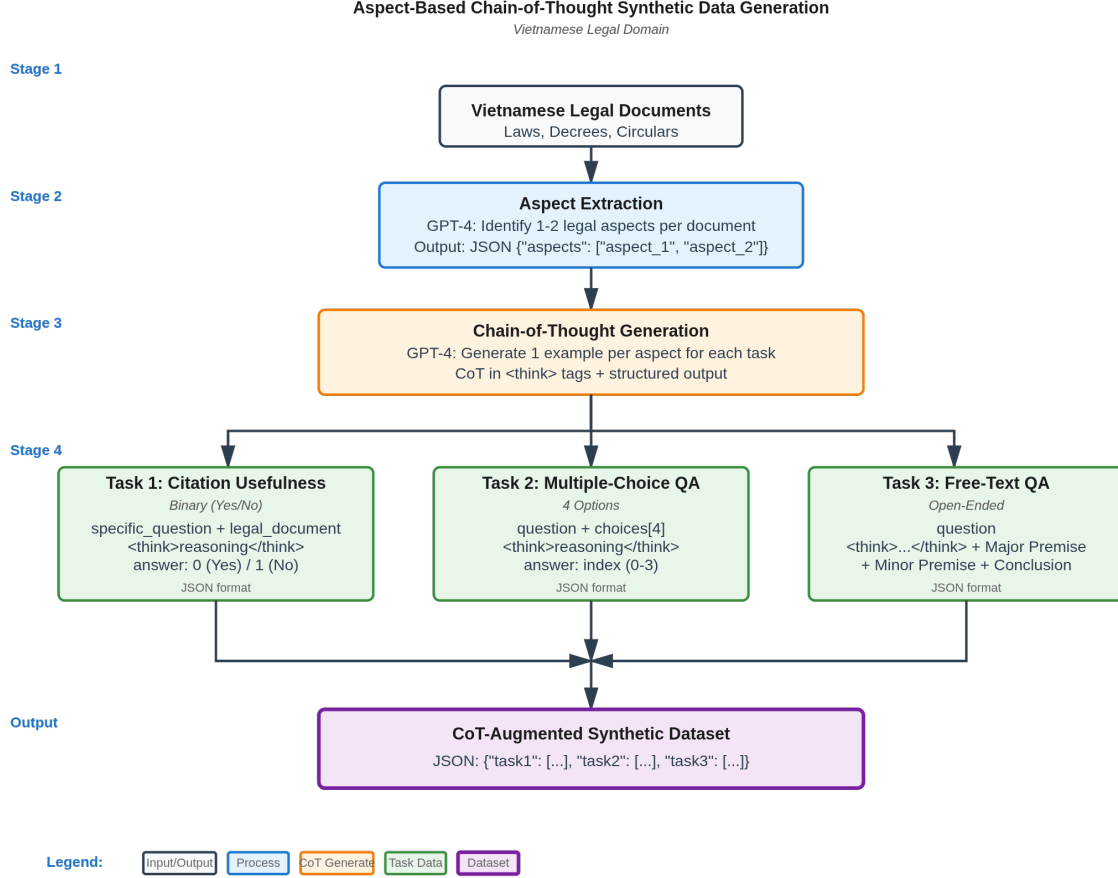


Figure 2: All the stages of synthetic data generation.

consistent dialogue structure across different legal domains; (3) response targets are carefully structured to include both final legal conclusions and intermediate reasoning chains, enabling the model to learn both the decision-making process and the final output generation.

QLoRA adapter design and quantization strategy: We adopt QLoRA (Dettmers et al., 2023) as our primary PEFT mechanism, which represents a sophisticated approach to large model adaptation through strategic combination of low-rank decomposition and advanced quantization techniques. QLoRA operates by injecting small, trainable low-rank adapter matrices into selected Transformer projection layers while maintaining the original model weights in a frozen, quantized state. This architectural choice enables efficient instruction tuning of models containing tens of billions of parameters on single GPU configurations while preserving the performance characteristics typically associated with full 16-bit precision fine-tuning.

The quantization component of QLoRA lever-

ages several advanced techniques: (1) 4-bit NormalFloat (NF4) quantization, which is specifically designed to handle the typical distribution patterns found in neural network weights, providing superior preservation of model capacity compared to uniform quantization schemes; (2) double quantization, an innovative approach that applies a second level of quantization to the quantization constants themselves, further reducing memory overhead without significant precision loss; (3) paged optimizers that utilize NVIDIA’s Unified Memory feature to automatically manage GPU-CPU memory transfers during training, preventing out-of-memory conditions during gradient computation and parameter updates.

4 Experiments

4.1 Experimental Settings

Table 3 outlines the main training setup. The use of 4-bit quantization and LoRA enabled memory-efficient fine-tuning of the 4B-parameter model, with only 1.62% of parameters updated. The

dataset combined reasoning and non-reasoning legal texts, while optimization with AdamW and a conservative learning rate ensured stable convergence. An effective batch size of 128 was achieved through gradient accumulation, allowing efficient training under hardware constraints.

4.2 Results and Discussion

Table 1 shows the performance of our approach compared with different variants on the public test set. The experimental results demonstrate a clear hierarchy in model performance based on training strategies and features. The Qwen3-finetune + DA model, when trained with the FPFT strategy and incorporating Aspect-based fine-tuning, consistently outperformed all other variants across both the QA and NLI tasks in both thinking and non-thinking modes. This configuration achieved the highest scores in three out of four metrics: a remarkable 89.73 on the Thinking Mode QA Task, 96.00 on the Thinking Mode NLI Task, and 91.10 on the Non-thinking Mode QA Task. The only metric where this model was not the top performer was the Non-thinking Mode NLI Task, where another Qwen3-finetune + DA variant with the QLoRA training strategy and Aspect-based features achieved a slightly higher score of 97.33. A significant finding is the impact of Data Augmentation, which led to a substantial improvement in performance for the QLoRA and FPFT models, especially in the more challenging Thinking Mode.

The observed results suggest that deeper fine-tuning and specialized training features are critical for maximizing model performance on complex legal tasks. The FPFT strategy consistently showed a marked advantage over QLoRA, particularly when combined with Data Augmentation, indicating that adjusting all model parameters is more effective than low-rank adaptation for these specific tasks. This superiority is likely due to the ability of FPFT to better capture the nuanced and intricate patterns inherent in legal data. Furthermore, the significant boost in performance seen with the Aspect-based fine-tuning highlights its crucial role. This suggests that explicitly training the model to consider specific aspects of legal documents and questions enhances its ability to reason and make more accurate predictions, especially within the more cognitively demanding Thinking Mode. In conclusion, the combination of a comprehensive fine-tuning approach like FPFT, enriched with data augmentation and aspect-based training, represents a robust

and effective strategy for achieving state-of-the-art performance in legal domain language model applications.

As shown in Table 2, we achieved the top rank on the final LegalSLM Leaderboard with the highest average score of 81.08. We outperformed the second-place team, MinLegal, by a notable margin, demonstrating our superior performance across the board. Our success was driven by a commanding lead in the QA task, where our score of 92.67 was the highest among all competitors. Additionally, our performance in NLI task was exceptionally strong, with a score of 97.00 placing us as the second-best in this category.

5 Error Analysis and Case Study

The quantitative enhancements are derived from the qualitative variations in the training data. The aspect-based generating method creates training examples that are more targeted and full of context. The synthetic data makes the model understand how different components of a legal rule are related by breaking down legal articles into particular parts, like “conditions”, “penalties”, “scope of application”, and “exceptions”. On the other hand, standard data generation may create more basic question-answer pairs that capture the main idea of an article but omit its more subtle details.

One clear example of this improvement involves a question about land allocation limits under the 2024 Land Law, as shown in Table 4, asking to compare the limit for perennial crops in the delta region versus the midland and mountainous regions.

- The **standard model** incorrectly answered that the allocation limit is higher in the delta. This is a factual error that suggests a superficial understanding, potentially confusing general economic productivity with specific legal allocations.
- The **aspect-based model**, however, correctly identified that the land allocation limit is higher in the midland and mountainous regions. Its reasoning directly cited the relevant article, demonstrating a precise understanding of how the law applies differently based on the "geographic area" aspect. This shows the model's enhanced ability to parse and apply rules that have conditional scopes.

Table 1: The experimental results on the public test.

Model Variants	Training Strategy	Thinking Mode		Non-thinking Mode		Aspect-based
		QA Task	NLI Task	QA Task	NLI Task	
Qwen3-base-legal	Baseline	71.23	84.00	82.19	92.67	No
Qwen3-finetune	QLoRA	68.49	91.33	84.59	96.67	No
Qwen3-finetune + DA	QLoRA	81.51	95.33	86.99	95.33	No
Qwen3-finetune + DA	FPFT	86.30	96.00	88.36	96.00	No
Qwen3-finetune + DA	QLoRA	85.62	91.33	85.62	97.33	Yes
Qwen3-finetune + DA	FPFT	89.73	96.00	91.10	96.00	Yes

Table 2: The final LegalSLM Leaderboard on the private test. The best scores are in bold, and the second-best scores are underlined.

Rank	Team name	vi-law-nli	vi-law-qa	vi-law-syllo	Average
2	MinLegal	98.00	<u>87.33</u>	53.08	<u>79.47</u>
3	URAx	94.50	83.33	57.67	78.50
4	Innovation-LLM	95.67	83.67	<u>54.17</u>	77.84
5	LICTU	84.67	80.67	53.75	73.03
1	Bosch@AI Team (Ours)	<u>97.00</u>	92.67	53.58	81.08

A second compelling example relates to the legal consequences of illegal drug use refer to table 5. When asked how this act is handled under current law, the models diverged significantly.

- The **standard model** made a critical legal error, stating that the act results in criminal prosecution. This conflates a civil administrative penalty with a much more severe criminal charge.
- The **aspect-based model** correctly answered that the act is an administrative violation subject to fines, not criminal prosecution. This distinction is fundamental in legal practice. The model’s success here strongly suggests that the aspect-based data, with its focus on the "penalty" aspect, trained it to differentiate between various types of legal sanctions and apply them to the correct context.

The NLI job also shows this increased nuance. Envision a scenario where an individual employs counterfeit documentation to register for a vehicle auction, anticipating disqualification.

- The conventional model correctly stated that the action was illegal, although it failed to provide a precise rationale, merely asserting that the use of fraudulent documents is unlawful. It failed to associate the precise conduct with the particular consequence of disqualification.

- Conversely, the aspect-based method provided a more clear rationale. It delineated the specific act (giving fraudulent information and documents) as a direct violation articulated within the "grounds for disqualification" provision of the pertinent decree. This signifies a more sophisticated understanding, transitioning from a generic concept of illegality to the application of a specific regulation to achieve a particular outcome.

This improved capability stems from the nature of the aspect-based data, which likely presented the model with scenarios where various forms of providing false information were explicitly linked to the "disqualification" penalty. This structured learning helps the model build a more accurate map of legal cause and effect.

Overall, the integration of aspect-based synthetic data generation marks a significant step forward in fine-tuning smaller language models for the specialized domain of Vietnamese law. The method not only yields superior performance on quantitative benchmarks for both multiple-choice and inference tasks but also, more importantly, fosters a deeper and more nuanced understanding of legal texts. By training the model on data that deconstructs legal articles into their core functional components, we enable it to move beyond keyword matching and surface-level comprehension. It learns to reason about legal principles, conditions, and consequences in a manner that more

closely mimics expert legal analysis. The specific examples of improved performance on complex questions highlight the value of this structured data approach in building more reliable and accurate legal AI systems.

6 Conclusion and Future Work

This paper presents our solution for LegalSML task of VLSP-2025 shared-task. We developed a framework that combines synthetic instruction tuning with chain-of-thought question answer data to guide a large teacher model. This data was then used with both Parameter-Efficient Fine-Tuning and Full-parameter Fine-Tuning to adapt the teacher model into a compact student model optimized for the Vietnamese legal domain. Our system achieved first place in average scores and demonstrated competitive performance across three evaluation tasks.

In the future, we will focus on improving performance in the free-text question answering task. In addition, we plan to optimize fine-tuning techniques for small language models specifically designed for the Vietnamese language.

References

- Tim Dettmers, Artidoro Pagnoni, Ari Holtzman, and Luke Zettlemoyer. 2023. Qlora: Efficient finetuning of quantized llms. *Advances in neural information processing systems*, 36:10088–10115.
- Edward J Hu, Yelong Shen, Phillip Wallis, Zeyuan Allen-Zhu, Yuanzhi Li, Shean Wang, Lu Wang, Weizhu Chen, and 1 others. 2022. Lora: Low-rank adaptation of large language models. *ICLR*, 1(2):3.
- Jean Lee, Nicholas Stevens, and Soyeon Caren Han. 2025. Large language models in finance (finllms). *Neural Computing and Applications*, pages 1–15.
- Chen Ling, Xujiang Zhao, Jiaying Lu, Chengyuan Deng, Can Zheng, Junxiang Wang, Tanmoy Chowdhury, Yun Li, Hejie Cui, Xuchao Zhang, and 1 others. 2023. Domain specialization as the key to make large language models disruptive: A comprehensive survey. *arXiv preprint arXiv:2305.18703*.
- Aixin Liu, Bei Feng, Bing Xue, Bingxuan Wang, Bochao Wu, Chengda Lu, Chenggang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, and 1 others. 2024. Deepseek-v3 technical report. *arXiv preprint arXiv:2412.19437*.
- Fenglin Liu, Hongjian Zhou, Boyang Gu, Xinyu Zou, Jinfa Huang, Jing Wu, Yiru Li, Sam S Chen, Yining Hua, Peilin Zhou, and 1 others. 2025. Application of large language models in medicine. *Nature Reviews Bioengineering*, pages 1–20.
- Henrique Marcos. 2025. Can large language models apply the law? *AI & SOCIETY*, 40(5):3605–3614.
- Humza Naveed, Asad Ullah Khan, Shi Qiu, Muhammad Saqib, Saeed Anwar, Muhammad Usman, Naveed Akhtar, Nick Barnes, and Ajmal Mian. 2025. A comprehensive overview of large language models. *ACM Trans. Intell. Syst. Technol.*, 16(5).
- Minh Thuan Nguyen, Khanh-Tung Tran, Nhu Van Nguyen, and Xuan-Son Vu. 2023. Vigptqa-state-of-the-art llms for vietnamese question answering: system overview, core models training, and evaluations. In *Proceedings of the 2023 conference on empirical methods in natural language processing: industry track*, pages 754–764.
- Tan-Minh Nguyen, Hoang-Trung Nguyen, Trong-Khoi Dao, Xuan-Hieu Phan, Ha-Thanh Nguyen, and Thi-Hai-Yen Vuong. 2025. Vlqa: The first comprehensive, large, and high-quality vietnamese dataset for legal question answering. *arXiv preprint arXiv:2507.19995*.
- Long Ouyang, Jeffrey Wu, Xu Jiang, Diogo Almeida, Carroll Wainwright, Pamela Mishkin, Chong Zhang, Sandhini Agarwal, Katarina Slama, Alex Ray, and 1 others. 2022. Training language models to follow instructions with human feedback. *Advances in neural information processing systems*, 35:27730–27744.
- Fali Wang, Minhua Lin, Yao Ma, Hui Liu, Qi He, Xianfeng Tang, Jiliang Tang, Jian Pei, and Suhang Wang. 2025. A survey on small language models in the era of large language models: Architecture, capabilities, and trustworthiness. In *Proceedings of the 31st ACM SIGKDD Conference on Knowledge Discovery and Data Mining V. 2*, pages 6173–6183.
- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, and 1 others. 2025. Qwen3 technical report. *arXiv preprint arXiv:2505.09388*.
- Qin Zhang, Ziqi Liu, and Shirui Pan. 2025. The rise of small language models. *IEEE Intelligent Systems*, 40(1):30–37.
- Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Siddhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou, and Le Hou. 2023. Instruction-following evaluation for large language models. *arXiv preprint arXiv:2311.07911*.

A Experimental Setting

We present the model configurations for fine-tuning LLM with the LoRA technique.

Table 3: Training configuration for fine-tuning Legal Qwen-4B on legal QA.

Component	Configuration
Base model	Legal Qwen-4B pretrained ¹
Model size	4B parameters
Quantization	4-bit (Unsloth)
Max sequence length	2048 tokens
Fine-tuning method	LoRA (parameter-efficient)
LoRA rank (r)	32
LoRA alpha	64
LoRA target modules	q_proj, k_proj, v_proj, o_proj, gate_proj, up_proj, down_proj
Trainable parameters	~66M (1.62% of total)
Dataset	Reasoning-focused + non-reasoning legal data (200K examples)
Epochs	1
Trainer	SFTTrainer (TRL library)
Optimizer	AdamW (8-bit)
Learning rate	2e-5 (linear scheduler, 5 warmup steps)
Batch size (per device)	16
Gradient accumulation steps	8
Effective batch size	128

¹ <https://huggingface.co/VLSP2025-LegalSML/qwen3-4b-legal-pretrain>

B Appendix: LLM Prompts Used in Experiments

B.1 Prompt 1: Aspect Extraction

Prompt 1: Aspect Extraction

Bạn là một chuyên gia pháp luật Việt Nam. Dựa trên văn bản pháp luật được cung cấp, hãy xác định 1-2 khía cạnh pháp lý riêng biệt, điểm chính hoặc chủ đề được bao quát.

Thông tin văn bản: {document context}

Nội dung văn bản: {doc_info["content"]}

Trả về **CHỈ MỘT** đối tượng JSON: { "aspects": ["Khía cạnh 1 ngắn gọn", "Khía cạnh 2", ...] }

B.2 Prompt 2: Reasoning Generation

Prompt 2: Reasoning Generation

Part 1: Reasoning System Prompt

Bạn là một chuyên gia pháp luật Việt Nam chuyên về tổng hợp dữ liệu.

Dựa trên văn bản pháp luật được cung cấp và các khía cạnh đã xác định, tạo ra 1 ví dụ tổng hợp cho mỗi khía cạnh cho mỗi loại trong 3 dạng nhiệm vụ sau.

Yêu cầu chung:

- Đảm bảo các ví dụ chính xác, đa dạng và dựa trực tiếp trên văn bản pháp luật được cung cấp.
- Sử dụng tiếng Việt cho tất cả câu hỏi, lựa chọn và câu trả lời.

Thông tin đầu vào:

- Thông tin văn bản: {document_context}
- Nội dung văn bản: {doc_info["content"]}
- Các khía cạnh pháp lý đã xác định: {aspects_list}

Part 2: Task 1 - Legal Citation Classification

Nhiệm vụ 1: Tính hữu ích của trích dẫn pháp luật (Phân loại Đúng/Sai)

- Cho mỗi khía cạnh, tạo ví dụ trong đó:
 - legal_document: một trích dẫn có liên quan từ văn bản đầu vào
 - specific_question: một câu hỏi pháp lý liên quan đến khía cạnh
 - question: luôn cố định là "Điều luật được cung cấp có thể dùng để trả lời câu hỏi trên hay không?"
 - choices: ["Có", "Không"]
 - answer: 0 cho "Có", 1 cho "Không"
- Xen kẽ ví dụ đúng (hữu ích) và sai (không liên quan), bắt đầu với đúng cho khía cạnh đầu tiên
- Reasoning: Giải thích ngắn gọn chi tiết về việc trích dẫn này có hoặc không hữu ích, bao bọc trong thẻ <think> ... </think>
- Định dạng: dict JSON

Part 3: Task 2 - Multiple Choice Questions

Nhiệm vụ 2: Câu hỏi trắc nghiệm pháp luật

- Cho mỗi khía cạnh, tạo câu hỏi kiểm tra kiến thức toàn diện từ văn bản đầu vào
- Cung cấp 4 lựa chọn (choices), với một answer đúng (chỉ số 0–3)
- Các lựa chọn gây nhiễu phải hợp lý nhưng không chính xác
- Reasoning: Giải thích chi tiết tại sao đáp án đúng và các lựa chọn khác sai, bao bọc trong thẻ <think> ... </think>
- Định dạng: dict JSON với trường question, choices (list), answer (int)

Part 4: Task 3 - Open-ended Legal Questions

Nhiệm vụ 3: Câu hỏi pháp luật dạng tự luận

- Cho mỗi khía cạnh, tạo câu hỏi mở đòi hỏi trả lời dạng tường thuật
- Cấu trúc câu trả lời:
 - Phần phân tích chi tiết nằm trong thẻ <think> ... </think>
 - Sau đó lần lượt ghi:
 - * Tiền đề lớn: ...
 - * Tiền đề nhỏ: ...
 - * Kết luận: ...
- Thứ tự khuyến nghị:
 1. Phân tích (trong <think>)
 2. Tiền đề lớn
 3. Tiền đề nhỏ
 4. Kết luận

Part 5: Output Format

Định dạng đầu ra: Trả về một đối tượng JSON với ba khóa: "task1", "task2", "task3", mỗi khóa chứa các mảng ví dụ.

C Appendix: Example of Instruction Dataset for LoRA Fine-tuning

C.1 An Example of Converting Instruction Data into Qwen Chat Format

Original Conversation Before Conversion

```
[
  {
    "role": "system",
    "content": "Bạn là một chuyên gia pháp luật Việt Nam.  
Nhiệm vụ của bạn là xác định xem một điều luật có thể  
được sử dụng để trả lời câu hỏi pháp lý cụ thể hay không."
  },
  {
    "role": "user",
    "content": "Điều luật: Điều 4, Nghị định 33/2024/NĐ-CP: 1. Vũ khí hoá học  
được định nghĩa theo Công ước cấm vũ khí hoá học bao gồm  
các loại sau đây, riêng biệt hoặc trong tổ hợp: ...  
  
Câu hỏi: Điều 4 có quy định về quy trình thanh tra hóa chất Bảng  
tại cơ sở sản xuất không?  
  
Điều luật được cung cấp có thể dùng để trả lời câu hỏi trên hay không?"
  },
  {
    "role": "assistant",
    "content": "<think>Điều 4 chỉ giải thích về khái niệm vũ khí hóa học,  
không quy định về quy trình thanh tra hóa chất Bảng tại cơ sở sản xuất.  
Vì vậy, điều luật này không hữu ích để trả lời câu hỏi trên.</think> Không"
  }
]
```

Formatted Conversation in Qwen Chat Format

```
<|im_start|>system
Bạn là một chuyên gia pháp luật Việt Nam.
Nhiệm vụ của bạn là xác định xem một điều luật có thể được sử dụng
để trả lời câu hỏi pháp lý cụ thể hay không.<|im_end|>

<|im_start|>user
Điều luật: Điều 4, Nghị định 33/2024/NĐ-CP:
1. Vũ khí hoá học được định nghĩa theo Công ước cấm vũ khí hoá học
bao gồm các loại sau đây, riêng biệt hoặc trong tổ hợp: ...

Câu hỏi: Điều 4 có quy định về quy trình thanh tra hóa chất Bảng
tại cơ sở sản xuất không?

Điều luật được cung cấp có thể dùng để trả lời câu hỏi trên hay không?<|im_end|>

<|im_start|>assistant
<think>
Điều 4 chỉ giải thích về khái niệm vũ khí hóa học, không quy định về
```

quy trình thanh tra hóa chất Bảng tại cơ sở sản xuất.
Vì vậy, điều luật này không hữu ích để trả lời câu hỏi trên.
</think>

Không<|im_end|>

Formatted Conversation in Qwen Chat Format and no Reasoning

<|im_start|>system
Bạn là một chuyên gia pháp luật Việt Nam.
Nhiệm vụ của bạn là xác định xem một điều luật có thể được sử dụng
để trả lời câu hỏi pháp lý cụ thể hay không.<|im_end|>

<|im_start|>user
Điều luật: Điều 4, Nghị định 33/2024/NĐ-CP:
1. Vũ khí hoá học được định nghĩa theo Công ước cấm vũ khí hoá học
bao gồm các loại sau đây, riêng biệt hoặc trong tổ hợp: ...

Câu hỏi: Điều 4 có quy định về quy trình thanh tra hóa chất Bảng
tại cơ sở sản xuất không?

Điều luật được cung cấp có thể dùng để trả lời câu hỏi trên hay không?<|im_end|>

<|im_start|>assistant
Không<|im_end|>

D Appendix: Qualitative Analysis: The Impact of Aspect-Based Data

Table 4: Land Allocation Limit Analysis (2024 Land Law)

Item	Content
Question	So sánh hạn mức giao đất trồng cây lâu năm cho cá nhân ở khu vực đồng bằng và khu vực trung du, miền núi theo quy định của Luật Đất đai 2024, điểm khác biệt chính là gì? (Compare the land allocation limit for perennial crops for individuals in the delta region versus the midland and mountainous regions according to the 2024 Land Law, what is the main difference?)
Correct Answer	Hạn mức giao đất ở trung du, miền núi cao hơn ở đồng bằng. (The land allocation limit is higher in the midland and mountainous regions than in the delta.)
Standard Model Prediction	Incorrect. Hạn mức giao đất ở đồng bằng cao hơn ở trung du, miền núi. (The land allocation limit is higher in the delta than in the midland and mountainous regions.)
Aspect-Based Model Prediction	Correct. Hạn mức giao đất ở trung du, miền núi cao hơn ở đồng bằng. (The land allocation limit is higher in the midland and mountainous regions than in the delta.)

Table 5: Legal Consequences of Drug Use

Item	Question / Response
Question	Hành vi sử dụng trái phép chất ma túy bị xử lý như thế nào theo quy định pháp luật hiện hành? (<i>How is the act of illegally using drugs handled under current law?</i>)
Correct Answer	Không bị truy cứu trách nhiệm hình sự nhưng bị xử phạt vi phạm hành chính... (<i>Not subject to criminal prosecution but subject to administrative sanctions...</i>)
Standard Model Prediction	Incorrect. Bị truy cứu trách nhiệm hình sự và phạt tiền... (<i>Subject to criminal prosecution and fines...</i>)
Aspect-Based Model Prediction	Correct. Không bị truy cứu trách nhiệm hình sự nhưng bị xử phạt vi phạm hành chính... (<i>Not subject to criminal prosecution but subject to administrative sanctions...</i>)

Table 6: Analysis of a Legal Entailment Scenario Regarding Vehicle Plate Auctions

Component	Content
Legal Document (Premise)	Căn cứ tại Điều 15 Nghị định 156/2024/NĐ-CP quy định về trường hợp truất quyền tham gia đấu giá biển số xe như sau: Các trường hợp bị truất quyền tham gia đấu giá gồm có cung cấp thông tin, tài liệu sai sự thật; sử dụng giấy tờ giả mạo để đăng ký tham gia đấu giá... (<i>Pursuant to Article 15 of Decree 156/2024/ND-CP, cases for disqualification from participating in a vehicle number plate auction include: providing false information or documents; using forged papers to register for the auction...</i>)
Specific Question (Hypothesis)	Một cá nhân sử dụng giấy tờ giả để đăng ký tham gia đấu giá biển số xe thì sẽ bị truất quyền tham gia. (<i>An individual who uses forged papers to register for a vehicle number plate auction will be disqualified.</i>)
Correct Answer	Entailment
Standard Model Response	Correct, but generic reasoning. The model correctly infers entailment but its reasoning is superficial: "Using forged documents is against the law." It doesn't cite the specific consequence mentioned in the premise.
Aspect-Based Model Response	Correct and specific reasoning. The model correctly infers entailment and provides a precise justification: "The act of using forged papers to register is explicitly listed as a reason for disqualification in the provided legal text." This demonstrates a deeper understanding of the cause-and-effect relationship defined by the legal rule.