

Automating Legal Concept Interpretation with LLMs: Retrieval, Generation, and Evaluation

Anonymous ACL submission

Abstract

Legal articles often include vague concepts to adapt to the ever-changing society. Providing detailed interpretations of these concepts is a critical task for legal practitioners, which requires meticulous and professional annotations by legal experts, admittedly time-consuming and expensive to collect at scale. In this paper, we introduce a novel retrieval-augmented generation framework, **ATRI**, for **AuT**omatically **R**etrieving relevant information from past judicial precedents and **I**nterpreting vague legal concepts. We further propose a new benchmark, Legal Concept Entailment, to automate the evaluation of generated concept interpretations without expert involvement. Automatic evaluations indicate that our generated interpretations can effectively assist large language models (LLMs) in understanding vague legal concepts. Multi-faceted evaluations by legal experts indicate that the quality of our concept interpretations is comparable to those written by human experts. Our work has strong implications for leveraging LLMs to support legal practitioners in interpreting vague legal concepts and beyond.

1 Introduction

When legislative bodies enact laws, in order to make relatively fixed legal texts more applicable to an ever-changing society, the legal texts often contain some vague (Endicott, 2000) and open-textured (Hart and Green, 2012) concepts. For example, in the Criminal Law of the People’s Republic of China, the article corresponding to the crime of Theft states: "...入户盗窃...的, 处三年以下有期徒刑..." ("Whoever ... enter a dwelling to steal ... , shall be sentenced to imprisonment of not more than 3 years, ..."). The term "dwelling" is a vague concept, and the article does not provide a clear definition of what kind of places are applicable to the concept of "dwelling". As shown in Figure 1, does a school dormitory apply to the

concept of "dwelling"? The article itself cannot provide a clear answer.

When interpreting a vague concept using doctrinal methods, legal professionals need library-based studies, reading extensive textbooks and past judicial precedents (Tiller and Cross, 2006; Yung-chin Su, 2024). This is admittedly a labor-intensive task with high time costs. Thus, manually drafting legal concept interpretations is always hard to scale or keep up-to-date, sometimes even biased (Farnsworth et al., 2011).

To alleviate the burden on human experts, previous studies have attempted to use LLMs to interpret legal concepts. Savelka et al. (2023) utilizes GPT-4 to interpret open-textured legal concepts from statutory articles based on valuable sentences from case law. Jiang et al. (2024) use LLMs to generate stories for legal concepts to assist in legal education. However, previous works have not yet: (1) effectively identified cases relevant to a given vague concept from numerous judicial cases; (2) extracted concept-focused information from relevant cases and used an LLM to summarize the legal interpretation accordingly; (3) proposed a benchmark to automatically evaluate the quality of legal concept interpretations without relying on legal experts.

In this paper, we propose a novel Retrieval-Augmented Generation (RAG) (Lewis et al., 2020; Guu et al., 2020) framework, **ATRI**, for **AuT**omatically **R**etrieving relevant information from cases and **I**nterpreting vague legal concepts. Our framework first adopts LLMs to retrieve cases that are relevant to the vague concept from a case database, and then extracts concept-relevant key information from these cases. Finally, we employ LLMs to summarize the extracted information and generate the concept interpretations. Furthermore, we propose a new automatic evaluation benchmark to automatically assess the quality of the interpretations by testing the extent to which the generated interpretations help LLMs better determine whether

this concept applies to an unseen case. Experiments show that our method can produce high quality interpretations comparable to human experts. Our contributions are as follows:

- We propose ATRI, a framework that utilizes LLMs to automatically retrieve relevant information from cases and generates interpretations for a given vague legal concept.
- We introduce a challenging task, Legal Concept Entailment, to automatically evaluate and compare the quality of legal concept interpretations.
- Automatic and human evaluations demonstrate that LLM-generated concept interpretations not only help LLMs understand vague concepts, but also achieve high quality comparable to that written by legal experts.

2 Related Works

Legal interpretation has been a longstanding challenge in the field of legal NLP (Nyarko and Sanga, 2022). Initially, rule-based methods (Waterman and Peterson, 1981; Paquin et al., 1991) provide users with tribunal decisions and doctrinal works to establish the meaning of open-textured legal concepts in specific contexts. With the advancement of deep learning, research (Šavelka and Ashley, 2021a,b) has used pretrained language models to retrieve sentences from legal cases that are useful for explaining legal concepts.

With the rapid progress of large language models, recent studies have also tried to use LLM to interpret legal texts. Jiang et al. (2024) use LLMs to generate stories to make the law more accessible for the public. However, the story-based explanation is not precise enough to help legal professionals, like lawyers or judges. Coan and Surden (2024) uses GPT to directly generate constitutional interpretation and Engel and Kruse (2024) further adds relevant cases to the input as references. These studies illustrate that using LLM to interpret legal concepts is possible, although they only evaluate one or two concepts, and whether their method could generalize to other concepts is uncertain. Savelka et al. (2023) proposes a general framework that could leverage previous judgments to generate legal concept interpretation. It proves that augmenting the LLM with human-extracted relevant judgments could improve the interpretation quality and eliminate the issue of hallucination. However, the

explanatory sentences it uses are manually selected from judgments, which is costly.

Different from previous work, we propose a fully automatic framework for legal concept interpretation. This framework leverages existing cases to generate interpretations for any concept without requiring any human involvement. Moreover, previous studies have predominantly relied on human evaluation to assess the quality of interpretations generated by large language models. To provide an objective and reproducible evaluation benchmark, we introduce an automatic evaluation task called Legal Concept Entailment and provide a corresponding dataset.

3 Preliminaries

The basic method of legal experts for writing legal interpretation involves extensively reading a large volume of previous legal cases, books, papers, and other materials related to specific legal articles. They then provide detailed interpretations of the specific applications of these articles, especially concerning the vague concepts within them. Such vague concepts are prevalent in legal articles, and their boundaries are not clear or well-defined. In this work, we focus on interpreting vague concepts within legal articles.

We introduce a challenging task, Vague Legal Concept Interpretation, which aims to provide interpretations for vague concepts in articles based on past cases. Formally, we define the task of Vague Legal Concept Interpretation as follows. Given a legal article a and a vague concept c within it, our task is to generate a legal interpretation e for the concept c , detailing the circumstances under which c applies or not.

4 Legal Concept Interpreter

In order to obtain interpretations of vague legal concepts automatically, we design a framework, ATRI. Following the method of legal experts, our framework summarizes the specific applications of the vague concept in judicial practice based on relevant case judgments. Specifically, our framework is composed of three parts (Figure 1): (1) **Retrieve**: Retrieve case judgments that mention the concept. (2) **Filter&Extract**: Select cases where the concept is analyzed in detail within the judgments, and extract the reasons for the determination of the concept in these cases. (3) **Interpret**: Use LLMs to generate the interpretation of the concept based on

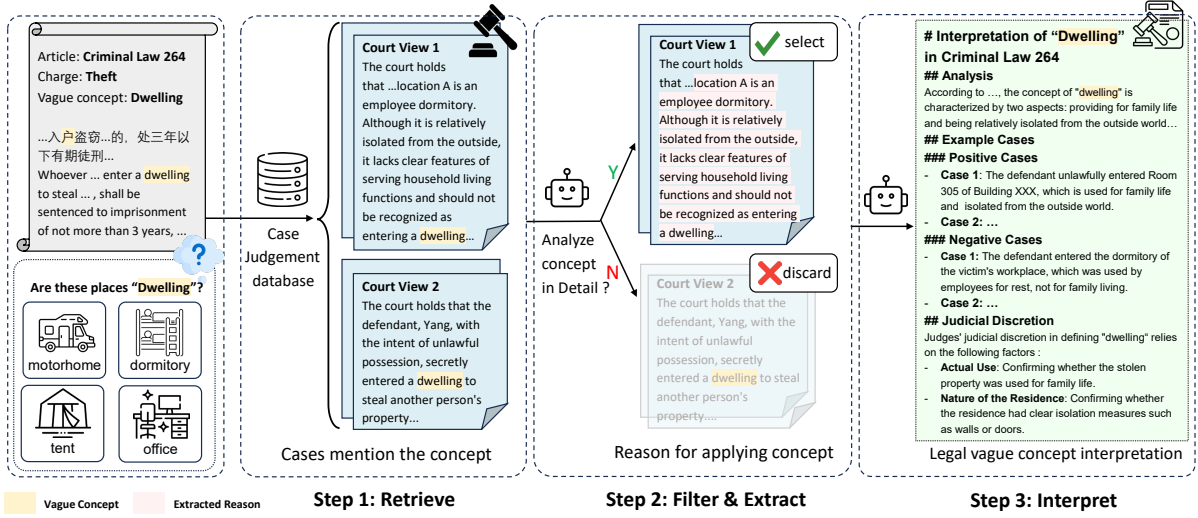


Figure 1: Overview of our framework, ATRI.

the extracted reasons.

4.1 Retrieving case judgments

To find case judgments that might be helpful to interpret the vague concept, the first step is to retrieve the cases that *mention* the concept. Formally, given a vague concept c and the article a that the concept belongs to, we will first find all the case judgments cited the Number of article a from a database that stores previous case judgments. Then we will retrieve the cases that mention the concept c through exact string matching and all of the retrieved cases form the set $\mathcal{D}_0$.

The case judgment database is constructed by collecting legal case judgments published on China Judgments Online¹. This is the largest public case judgment platform in China, which is the official website hosted by the Supreme People’s Court of China. Our database includes information from the years 1985 to 2021 available on the website, ensuring the comprehensiveness of the source.

A case judgment typically contains 5 parts: Header, Facts, Court view, Verdict and Conclusion². Among them, the court view section explains the legal rationale and basis for the judgment. We adopt a retrieval approach based on exact string matching to check whether the vague concept appears in the court view section of a case judgment. We do not use dense retrieval or other fuzzy matching methods for retrieving because legal terminology is very rigorous. Every concept has a fixed expression and rarely has alternative formulations.

¹<https://wenshu.court.gov.cn/>

²The details of the case judgment structure are provided in Appendix A.

Therefore, we use exact string matching to ensure retrieval precision.

4.2 Filtering relevant case judgments and Extracting reasons

In this step, we aim to filter relevant cases from the cases that just mention the concept and extract reasons for the determination of the vague concept in these cases. We define *relevant cases* of a concept as those cases in which the court view section provides a detailed reason why the vague concept applies to the case or not. We want to filter relevant cases because not all judgments that mention a concept are valuable for generating the interpretation of the concept. Some cases are relatively simple or straightforward, and judges may not provide detailed discussions of the concept in the judgment³.

We first use LLMs to filter the relevant cases from $\mathcal{D}_0$. Taking the court view as input, we first require the LLM to determine whether it provides a detailed reason, r , for why the concept c applies or not. If so, then extract this reason. The reason r should be a combination of original sentences from the court view. Next, we prompt the LLM to determine whether the concept applies to the case based on the court view, yielding a binary label l (Yes/No). The prompt we use for filtering is shown in Appendix F. From this process, we obtain a refined case set $\mathcal{D}_1$ containing cases that discuss the concept in detail in the court view.

Upon analyzing the labels within $\mathcal{D}_1$, we observe the proportion of positive cases (where c applies to

³We show an example of a judgment that mentions the concept only and a relevant case judgment that discusses the concept in detail in Appendix B

the case) far exceeds negative cases, with a ratio surpassing 100:1. This may be because the cases we collected are all prosecuted and adjudicated cases. In judicial practice, only cases where substantial evidence supports prosecution are brought to court, making it more likely that the concept applies to these cases, resulting in a higher proportion of positive examples. To comprehensively account for different situations when generating concept interpretations, we aim to ensure that both positive and negative examples receive adequate attention. Therefore, we only sample a subset of positive cases to construct a balanced dataset, $\mathcal{D}$, and its corresponding reason set $\mathcal{R}$.

4.3 Generating Concept Interpretations

After collecting the relevant cases and reasons, this step leverages a LLM to summarize these past experiences and generate an interpretation of the vague concept.

An interpretation should elaborate on how the vague concept has been explained or applied by the courts. We designed the interpretation to consist of three main components: *Analysis*, which explains the basic meaning of the concept and its applicability conditions; *Case Examples*, which provide representative positive and negative cases from past rulings; and *Judicial Discretion*, which offers criteria to guide judges in flexibly applying vague concepts based on case specifics. (see Appendix E)

The input to the LLM for generating interpretations consists of the following components: (1) legal article a , (2) vague concept c (3) reason set $\mathcal{R}$ (4) interpretation example e_0 . We require the output interpretation to follow the same format as the interpretation example e_0 , to ensure a consistent and standardized format. (see Appendix E.2)

5 Is generated interpretation reliable?

To evaluate the quality of the generated interpretations, previous work has predominantly relied on human evaluation. We also conducted a human evaluation, as detailed in Section 7. However, human evaluation is inherently subjective, and we aim to assess the quality of the generated concepts in a more objective and quantitative manner. Therefore, we propose a new benchmark, *Legal Concept Entailment*, which is reproducible and enables an objective comparison of interpretations generated by different methods.

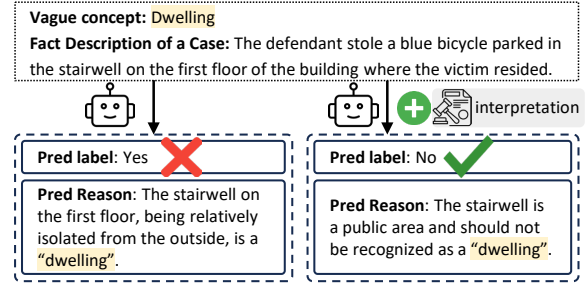


Figure 2: An example of Legal Concept Entailment Task. The left half of the figure illustrates the LLM directly performing the task, while the right half shows the LLM completing the task with the concept interpretation as a reference.

5.1 Legal Concept Entailment Task

If an interpretation of a concept is good, it should assist humans or models in better determining whether the concept applies to an unseen case and in providing the corresponding reasoning. Based on this assumption, we design the *Legal Concept Entailment* task. Given the fact description of a case relevant to the vague concept, the task is to determine whether the concept applies and provide a reason. We use a fixed LLM, to perform this classification task. By incorporating interpretations from different sources into the input, we can observe changes in the classification accuracy, which allows us to assess the quality of the interpretations. The more accurate the classification, the higher the quality of the interpretation.

Formally, this task is divided into two parts. The first part is a binary classification task: for a vague concept c in a legal article a , given the fact description f of an unseen relevant case d , the output should be a binary label $\hat{l}$ (Yes/No), indicating whether c applies to the fact f . The second part is a generation task, which requires generating a reason $\hat{r}$ to explain the prediction result of the binary classification task. An example is shown in Fig 2.

5.2 Dataset

We recruited a legal expert with substantial judicial experience to identify frequently encountered legal articles and vague concepts in judicial practice. Specifically, we selected articles cited in the case database by more than 10,000 cases, resulting in a total of 14 articles and 16 concepts.

For each concept, we collected relevant cases, which are those where judges provided detailed discussions of the concept in their rulings. These cases are considered challenging, thus warranting such

detailed deliberation. To gather this data, we reused the retrieval and filtering modules described in Section 4.1 and Section 4.2. On average, 166 cases were selected for each concept, with a positive-to-negative case ratio of 2:1. Detailed statistics are provided in Appendix G.

Following methods outlined in Sec 4.2, we use Qwen2.5-72B-Instruct (Qwen Team, 2024) to annotate each case with the gold label l and reason r for the Legal Concept Entailment task. Manual inspection indicates that the annotated data is highly accurate (see Appendix C.1).

The distinction between data annotation and the Legal Concept Entailment task lies in the input provided to the LLM. For annotation, the input is the court view, which contains explicit judgments made by judges and can be directly extracted as ground truth. In contrast, for the task itself, the input is the fact description, which lacks explicit judgments, requiring the LLM to perform reasoning to infer the entailment.

5.3 Evaluation Metrics

For the classification task, we use Accuracy (Acc.), Macro Precision (Ma-P), Macro Recall (Ma-R), and Macro F1 (Ma-F) as the evaluation metrics. The reason for using macro average is that the number of cases relevant to each concept is imbalanced, and we want to give equal weight to all classes.

For the reason generation task, we use a GPT-4-based evaluator to evaluate the consistency between the generated reason $\hat{r}$ and the gold reason r from the court view, following previous LLM-as-a-Judge based methods (Zheng et al., 2023; Zhu et al., 2023). We require the GPT-4 to rate from 1 to 10 for the consistency between the $\hat{r}$ and r , with higher scores indicating greater consistency. Specifically, we use gpt-4o-2024-08-06 (Achiam et al., 2023) and set the temperature to 0. The prompt we use for evaluation is in Appendix F.1.5. Note that, if the classification result is incorrect, the consistency score is directly set to 0.

5.4 Method

This section introduces how to perform the Legal Concept Entailment task with the incorporation of concept interpretations. First, we generate the interpretations for the concepts following the method described in Section 4. To prevent data leakage, the cases used for generating interpretations do not overlap with the test set. Next, we prompt the LLM

to perform the Legal Concept Entailment task using the generated interpretations (see Appendix F.1.6).

As shown in the right half of Figure 2, given a vague concept c in a legal article a and the fact description f of a relevant case d , the LLM is prompted to analyze whether the concept c applies to the fact f based on the concept interpretation. Specifically, the LLM first generates a reason $\hat{r}$ and subsequently assigns a classification label $\hat{l}$.⁴

5.5 Baselines

We compare our method with two baseline categories: "w/o Interpretation," where the LLM relies solely on its internal knowledge for the Legal Concept Entailment Task, and "w/ Interpretation," where the LLM is provided with an interpretation of the vague concept for the task.

w/o Interpretation (1) **Random**: We use random guessing of "Yes" or "No" as a weak baseline. (2) **Zero-shot (ZS)**: The LLM performs the Legal Concept Entailment task in a zero-shot setting. Specifically, only the legal article a , the vague concept c , and the fact description f of the relevant case d are provided as input. (Shown in the left half of Figure 2.) (3) **Chain-of-Thought (Kojima et al., 2022)**: Using the prompt "Let's think step by step" to encourage the LLM to generate intermediate steps and improve its reasoning.

w/ Interpretation We introduced concept interpretations generated by different approaches, including human-written and model-generated interpretations, to compare them with our method: (1) **Judicial Interpretation (JI)**: We recruit a legal expert to retrieve judicial interpretations for the concept c . Judicial interpretations are explanations issued by the Supreme People's Court on how to specifically apply the law. (2) **Expert interpretation (EI)**: We collect legal professionals' interpretations for the concept c from FaXin⁵ and WeChat official accounts of major law firms, which are of high quality. (3) **LLM Direct Interpretation (DI)**: Without providing relevant cases, the LLM generates an interpretation of the vague concept c directly based on its internal knowledge.

5.6 Result

We report the performance of our method and all baselines on Legal Concept Entailment Task in

⁴Implementation details can be found in Appendix C.1

⁵<https://www.faxin.cn/>,

	Qwen2.5 (72B)					Qwen2.5 (14B)				
	Acc	Ma-P	Ma-R	Ma-F	CS	Acc	Ma-P	Ma-R	Ma-F	CS
Random	51.66	51.13	51.23	50.32	/	51.66	51.13	51.23	50.32	/
Zero-Shot	71.38	<u>72.64</u>	61.81	61.42	5.658	70.92	73.04	60.78	59.88	5.525
Chain-of-Thought	71.95	<u>72.07</u>	63.26	63.46	<u>5.717</u>	71.52	73.83	61.60	61.01	5.666
Judicial Interpretation	72.10	69.87	65.82	66.54	5.573	70.92	68.24	64.62	65.23	5.347
Expert interpretation	72.13	70.78	64.68	65.30	5.630	71.95	69.85	65.31	66.01	5.581
Direct Interpretation	<u>72.35</u>	70.03	<u>66.43</u>	<u>67.18</u>	5.642	<u>72.72</u>	70.98	<u>66.11</u>	<u>66.90</u>	<u>5.677</u>
ATRI (Ours)	75.03	73.21	69.97	70.87	5.946	74.50	<u>72.49</u>	69.56	70.39	5.840

Table 1: Main results of automatic evaluation on the Legal Concept Entailment Task, the best is **bolded** and the second is underlined. The metric is accuracy (Acc), Macro F1-score (Ma-F), and consistency score (CS).

Table 1. Overall, our ATRI achieves the best performance across nearly all models and evaluation metrics, showcasing the effectiveness of the proposed framework and the necessity of its core components.

5.6.1 Classification Task

For the classification task, we found that:

(1) LLMs possess some level of discriminative ability. The performance of "w/o Interpretation" surpasses that of random guessing. Besides, the performance of CoT surpasses that of Zero-shot, demonstrating that step-by-step reasoning is beneficial for the Legal Concept Entailment Task.

(2) Interpretations for vague concepts are valuable. The performance of "w/ Interpretation" significantly outperforms that of "w/o Interpretation". "w/ Direct Interpretation" shows that LLMs can leverage their extensive internal knowledge to reason about vague concepts and generate useful legal concept interpretations. "w/ Judicial Interpretation" falls short of "w/ Direct Interpretation". We attribute this to the relatively simple explanations provided in judicial interpretations, which lack the depth required to guide LLMs in evaluating the applicability of vague concepts to specific cases. The performance of "w/ Expert Interpretation" is inferior to ATRI. We attribute this to the fact that expert-written interpretations are often overly abstract and detailed, which results in poorer readability. We will discuss this in detail during the human evaluation (Sec 7).

(3) Utilizing relevant cases is necessary. ATRI outperforms "w/ Direct Interpretation", demonstrating the effectiveness of generating interpretations with reference to relevant cases.

5.6.2 Reason Generation Task

For the reason generation task, we found that: (1) The consistency score of ATRI is the highest, show-

ing significant improvement over both "w/o Interpretation" and "w/ Interpretation." baselines. This indicates that the interpretations generated by our method help the model better understand the concepts and make inferences in a correct manner. (2) Other "w/ Interpretation" methods generally perform worse than CoT, despite showing improvements in classification tasks. We believe this is because these interpretations are either incomplete or contain irrelevant information, which misguides the LLM to reason in the wrong direction.

5.7 Case study

Figure 3 presents an example of different methods applied to the Legal Concept Entailment Task. As demonstrated in the case, our framework accurately understands the applicability conditions of "dwelling" and outputs the right prediction with the right reasoning path. Our framework precisely captures the features of dwelling theft. In contrast, Zero-shot gave the wrong answer with the misunderstanding of the concept "dwelling". For "w/ Direct Interpretation", although it reaches the correct conclusion, the reasoning process contains errors and uncertainties. It failed to clarify the vague concept in this specific scenario, using the expression of "may not fully satisfy".

6 What affect interpretation quality?

In this section, we will discuss the impact of different settings in our framework ATRI on the quality of generated interpretations.

6.1 How to retrieve cases

To verify that using LLM to filter cases and perform label balancing on the case set can yield more useful relevant cases, we conducted the following experiments: (1) **No retrieval**, where no cases are retrieved (i.e. LLM Direct Interpretation); (2) **w/o Filtering**, which does not use LLM to filter cases or perform label balancing; (3) **w/o Balancing**, which

Fact Description of a Case: Defendant A went to the house of prostitute B and engaged in a sexual transaction with B. After the sexual transaction, defendant A took advantage of the opportunity to steal the property in B's house, which is equal to xxx RMB...		
Article 224: Whoever ... enter a dwelling to steal ... , shall be sentenced to imprisonment of not more than 3 years, ...		
Vague Concept: Dwelling	Charge: Theft	Gold Label: No
Gold Reason: The defendant entered the room in the name of prostitution, with the permission of the victim. At this time the room should be recognized as prostitution place, not for the victim's family life and relatively isolated from the outside world , should not be recognized as "dwelling".		
Zero-Shot	w/ Direct Interpretation	w/ Our Interpretation
Reason: Even if B's residence was used for illegal transaction, its nature is still as a private living space. Defendant A showed that his intent to illegally take possession of other people's property arose after he entered B's residence, and that the behavior did take place in B's private living space, posing a serious threat to B's personal safety and residential tranquility.	Reason: B's residence was mainly used for commercial activities (prostitution), which is different from the nature of ordinary family life. The residence may not fully satisfy the condition of "relative isolation from the outside world", as it has a certain degree of openness as a place of prostitution that frequently receives outsiders.	Reason: Although B's residence is somewhat enclosed and separated, its primary use is for sex transactional sexual activities rather than for family life . It did not contain the basic facilities and characteristics of everyday family. And the defendant entered the residence by way of prior agreement with B, rather than unlawful sneaking .
Label: Yes ✗	Label: No ✓	Label: No ✓

Figure 3: Case study for Legal Concept Entailment Task. The **green** parts indicate reasoning processes consistent with the gold reason, while the **red** parts indicate inconsistencies.

	Qwen2.5 (14B)	
	Ma-F	CS
No Retrieval	66.90	5.677
w/o Filtering	69.04	5.772
w/o Balancing	69.60	5.817
ATRI (Ours)	70.39	5.840

Table 2: Ablation study for relevant case retrieval method.

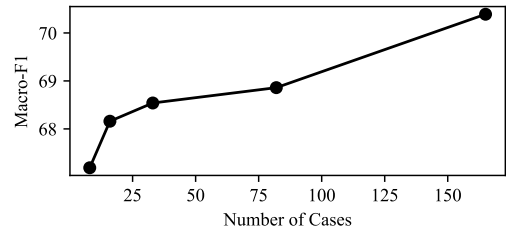


Figure 4: Results of different number of cases utilized to generate the interpretations. The model for generating interpretations and the model for prediction are Qwen2.5-72B and Qwen2.5-14B, respectively.

does not perform label balancing; We ensure that the number of cases retrieved by each method is consistent. The results in Table 2 indicate that every component of our retrieval method is useful for generating concept interpretations.

6.2 Number of cases

We investigated the impact of using different numbers of case judgments on the quality of generated concept interpretations. Specifically, we sampled different numbers of reasons from the extracted reason set $\mathcal{R}$ as input to the LLM. The results in Figure 4 demonstrate that a greater number of input reasons leads to higher-quality interpretations.

The more cases legal practitioners review, the more comprehensive their concept interpretations become. Our findings align with this process, showcasing LLMs' ability to analyze numerous cases effectively, highlighting their advantage to assist in legal concept interpretation.

6.3 Which parts of a case is useful?

In the second step of our framework, we only extract a few sentences discussing the concept from the court view of each relevant case, without including the complete fact description and court view. We aim to investigate whether this might result in the loss of important information from the case,

which could potentially affect the generation of interpretations. To explore this, we compared three different approaches to representing the information of a case during the interpretation generation step: (1) **Court view**: the part of the judgment where the judge explains the legal rationale and interprets the basis of the ruling; (2) **Summarized fact and Court view**: The facts section in case judgments is often lengthy and contains excessive detail. To address this, we first use an LLM to summarize the facts and then concatenate it with the court view section; (3) **Extracted reason from the court view**: Extracted reasons in Section 4.2.

	Qwen2.5 (14B)	
	Ma-F	CS
Court View	69.10	5.775
Fact & court view	70.17	5.818
Extracted Reason (Ours)	70.39	5.840

Table 3: Results of using different parts of case judgment to generate interpretations.

In the experiment, we control the number of input cases to be the same. In practice, using the "Extracted Reason" allows for the inclusion of more cases, as each entry is shorter in length. Even in this scenario with the same number of cases, we

see from Table 3 that "Extracted Reason" performs the best, indicating that it retains the important information while filtering out redundant details.

6.4 Components of Interpretation

In interpretation generation step, we ask the model to output the following components: Analysis, Example Cases, and Judicial Discretion. In this section, we will investigate whether each of these components is necessary. Specifically, we delete one main component at a time while keeping the other parts unchanged. The specific role of each component is detailed in Appendix E.

The results (Table 4) show that each component of the generated concept interpretation contributes to the overall performance. Notably, removing the "Example Cases" section results in the most significant performance drop, highlighting the importance of providing specific case examples.

	Qwen2.5 (14B) Macro-F1
w/o Example Cases	67.41
- w/o Positive Cases	68.17
- w/o Negative Cases	69.98
w/o Analysis	70.43
w/o Judicial Discretion	70.69
ATRI (Ours)	70.87

Table 4: Results of ablation experiments on different components of generated concept interpretations.

7 Why our interpretation better than others?

In the previous sections, we validated the effectiveness of our framework through performance on the Legal Concept Entailment Task. In this section, we further analyze the strengths of the generated interpretations through human evaluation.

7.1 Evaluation Metrics

We recruited 2 human evaluators with a legal education background to assess the legal concept interpretations generated by Qwen2.5 (72B). To provide a comprehensive evaluation of the interpretation, human raters judge the following metrics: (1) **Accuracy (Acc.)**, (2) **Informativeness (Info.)**, (3) **Normativity (Norm.)**, (4) **Comprehensiveness (Comp.)**, (5) **Readability (Read.)**. We use a 10-point Likert scale, where 1 represents "very poor" and 10 represents "very good".⁶

⁶Details about the metrics and human evaluation are discussed in Appendix D.

7.2 Results

In the top half of the Table 5, we have several interesting observations: (1) The average score of ATRI is the highest, indicating that ATRI can generate legal concept interpretations that are comparable to those produced by legal experts. (2) The Comprehensiveness score of ATRI is much higher than that of Expert Interpretation, indicating that ATRI, which involves having LLMs read a vast number of cases, effectively generates more comprehensive concept interpretations. (3) Expert Interpretation (EI) receives the lowest score in Readability, indicating that the interpretations written by legal experts tend to be abstract or complex, which hinders understanding by both humans and LLMs. (4) In Accuracy, Informativeness, and Normativity, ATRI shows improvements over Direct Interpretation (DI). Although there is still some gap with Expert Interpretation, it is important to note that Expert Interpretation was produced by legal experts who spent a considerable amount of time, while ATRI is already approaching human-level performance. In the future, combining the two approaches may be a better option, such as having the LLM first generate a draft, which can then be revised by legal experts to significantly improve efficiency.

	Acc.	Info.	Norm.	Comp.	Read.	Avg.
DI	7.03	6.21	7.53	<u>6.72</u>	7.38	6.97
EI	7.68	7.03	8.00	6.12	6.26	7.02
ATRI	<u>7.18</u>	<u>6.76</u>	<u>7.76</u>	7.15	<u>7.18</u>	7.21

Table 5: Human evaluation results of vague concept interpretations. The scores range from 1 to 10, with 10 being the highest. "Avg." represents the average score across five evaluation metrics.

8 Conclusion

In this work, we explore the novel application of LLMs in interpreting vague legal concepts. We propose a data-driven, fully automated framework to generate legal interpretations for a given vague concept. We further introduce a new task, Legal Concept Entailment, to automatically evaluate the generated interpretations. Both automatic and human evaluations demonstrate that our generated interpretations are useful and comparable to those written by legal experts. Our study suggests considerable potential for using LLMs in advancing legal interpretation and beyond.

Limitations

Sample Size Given the limited financial budget available to conduct our research, we chose to conduct our study on a smaller dataset to reduce the costs associated with using GPT-4 for scoring and hiring human evaluators. We would like to emphasize that even at this scale, the costs are not negligible. For example, evaluating the consistency score of the reasons across all experiments cost approximately \$300. For human evaluation, scoring the interpretations generated by the LLM and those written by legal experts cost around \$130.

Potential Risk of Dataset Leakage Although the large language model used in our experiments on Legal Concept Entailment Task is open-source, its training dataset is not fully transparent, which raises the possibility of data leakage. To address this issue, we evaluated various methods on the same base model to ensure a fair comparison. The relative improvements under different settings demonstrate our advantages.

Ethical Considerations

Privacy and Data Security Legal datasets frequently contain sensitive details about individuals and organizations, and improper handling can result in significant privacy violations. To safeguard this information, the case judgment dataset used in our experiments is thoroughly anonymized.

LLM-related Risks Large language models (LLMs) can inherit biases or inaccuracies from the data they are trained on, potentially leading to flawed legal interpretations. While LLMs can assist in generating legal concepts, they should not replace human judges or be used in real-world decision-making. Human oversight is essential to ensure fairness and accuracy in legal processes.

Code of Conduct This research follows the ACL Code of Ethics and respects participants' anonymity. We recruited two senior law school students for manual annotation and experiments, obtaining the participants' consent. We pay them wages higher than the local average hourly rate. We ensure that the content generated by the LLM was safe and non-offensive.

References

Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman,

Diogo Almeida, Janko Altenschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.

Andrew Coan and Harry Surden. 2024. Artificial intelligence and constitutional interpretation. *Arizona Legal Studies Discussion Paper*, (24-30).

Timothy AO Endicott. 2000. *Vagueness in law*. Oxford University Press.

Christoph Engel and Johannes Kruse. 2024. Professor gpt: Having a large language model write a commentary on freedom of assembly.

W. W. Farnsworth, Dustin F. Guzier, and Anup Malani. 2011. *Implicit bias in legal interpretation*.

Team GLM, Aohan Zeng, Bin Xu, Bowen Wang, Chenhui Zhang, Da Yin, Diego Rojas, Guanyu Feng, Hanlin Zhao, Hanyu Lai, Hao Yu, Hongning Wang, Jidai Sun, Jiajie Zhang, Jiale Cheng, Jiayi Gui, Jie Tang, Jing Zhang, Juanzi Li, Lei Zhao, Lindong Wu, Lucen Zhong, Mingdao Liu, Minlie Huang, Peng Zhang, Qinkai Zheng, Rui Lu, Shuaiqi Duan, Shudan Zhang, Shulin Cao, Shuxun Yang, Weng Lam Tam, Wenyi Zhao, Xiao Liu, Xiao Xia, Xiaohan Zhang, Xiaotao Gu, Xin Lv, Xinghan Liu, Xinyi Liu, Xinyue Yang, Xixuan Song, Xunkai Zhang, Yifan An, Yifan Xu, Yilin Niu, Yuantao Yang, Yueyan Li, Yushi Bai, Yuxiao Dong, Zehan Qi, Zhaoyu Wang, Zhen Yang, Zhengxiao Du, Zhenyu Hou, and Zihan Wang. 2024. *Chatglm: A family of large language models from glm-130b to glm-4 all tools*. *Preprint*, arXiv:2406.12793.

Kelvin Guu, Kenton Lee, Zora Tung, Panupong Pasupat, and Mingwei Chang. 2020. Retrieval augmented language model pre-training. In *International conference on machine learning*, pages 3929–3938. PMLR.

Herbert Lionel Adolphus Hart and Leslie Green. 2012. *The concept of law*. oxford university press.

Hang Jiang, Xiajie Zhang, Robert Mahari, Daniel Kessler, Eric Ma, Tal August, Irene Li, Alex Pentland, Yoon Kim, Deb Roy, and Jad Kabbara. 2024. *Leveraging large language models for learning complex legal concepts through storytelling*. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 7194–7219, Bangkok, Thailand. Association for Computational Linguistics.

Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. 2022. Large language models are zero-shot reasoners. *Advances in neural information processing systems*, 35:22199–22213.

Patrick Lewis, Ethan Perez, Aleksandra Piktus, Fabio Petroni, Vladimir Karpukhin, Naman Goyal, Heinrich Küttler, Mike Lewis, Wen-tau Yih, Tim Rocktäschel, et al. 2020. Retrieval-augmented generation for knowledge-intensive nlp tasks. *Advances in Neural Information Processing Systems*, 33:9459–9474.

- Julian Nyarko and Sarath Sanga. 2022. A statistical test for legal interpretation: Theory and applications. *The Journal of Law, Economics, and Organization*, 38(2):539–569.
- Louis-Claude Paquin, François Blanchard, and Claude Thomasset. 1991. Loge-expert: from a legal expert system to an information system for non-lawyers. In *Proceedings of the 3rd international conference on Artificial intelligence and law*, pages 254–259.
- Qwen Team. 2024. [Qwen2.5: A party of foundation models](#).
- Jaromír Šavelka and Kevin D Ashley. 2021a. Discovering explanatory sentences in legal case decisions using pre-trained language models. In *Findings of the Association for Computational Linguistics: EMNLP 2021*, pages 4273–4283.
- Jaromír Šavelka and Kevin D Ashley. 2021b. Legal information retrieval for understanding statutory terms. *Artificial Intelligence and Law*, pages 1–45.
- Jaromir Savelka, Kevin D Ashley, Morgan A Gray, Hannes Westermann, and Huihui Xu. 2023. Explaining legal concepts with augmented large language models (gpt-4). *arXiv preprint arXiv:2306.09525*.
- Emerson H Tiller and Frank B Cross. 2006. What is legal doctrine. *Nw. UL Rev.*, 100:517.
- DA Waterman and MA Peterson. 1981. Models of legal decision-making, r-2717-icj.
- Zhang Cheng Yung-chin Su, Zhou Xiang. 2024. The future of the legal dogmatics in the perspective of structure and management:new frontier of theoretical dogmatics with practical dogmatics in ai’s hand (in chinese). *NanJing University Law Journal*, 1:1–17.
- Lianmin Zheng, Wei-Lin Chiang, Ying Sheng, Siyuan Zhuang, Zhanghao Wu, Yonghao Zhuang, Zi Lin, Zhuohan Li, Dacheng Li, Eric Xing, et al. 2023. Judging llm-as-a-judge with mt-bench and chatbot arena. *Advances in Neural Information Processing Systems*, 36:46595–46623.
- Lianghui Zhu, Xinggang Wang, and Xinlong Wang. 2023. Judgelm: Fine-tuned large language models are scalable judges. *arXiv preprint arXiv:2310.17631*.

A The Structure of Case Judgments

A Case Judgment in China can generally be divided into five sections: the header, facts, court view, verdict, and conclusion. The **header** includes the name of the court, the type of document, case number, basic information about the parties involved, the origin of the case, and details about the judicial panel and trial method. The **facts** section outlines the plaintiff’s claims, facts, and arguments, as well as the defendant’s admissions regarding the plaintiff’s factual assertions. The **court view** section provides the rationale for the judgment and the legal basis upon which it is made. The **verdict** contains the decision on substantive issues of the case. Finally, the **conclusion** serves to formally end the judgment document.

B Examples of Relevant Cases

Charge	Vague concept	Cases mentioning the concept (Irrelevant Cases)	Cases that analyze the concept in detail (Relevant Cases)
Theft	Dwelling	The court holds that the defendant, Yang, with the intent of unlawful possession, secretly entered a <i>dwelling</i> to steal another person’s property. His actions constitute the crime of theft...	Regarding whether Zhang’s actions constitute theft by entering a <i>dwelling</i> , upon investigation, location A is an employee dormitory rented by B restaurant. Although it is relatively isolated from the outside, it lacks clear features of serving household living functions and should not be recognized as entering a <i>dwelling</i> .
Traffic accident crime	Flee the scene	After the accident, the defendant <i>fled the scene</i> and is fully responsible for the incident. His actions constitute the crime of traffic accident liability as stipulated in Article 133 of the Criminal Law of the People’s Republic of China.	The defendant argues that after the accident, he had his wife promptly dial 120 for emergency assistance and then left the scene to return home, claiming that he did not flee. Upon investigation, it is confirmed that the defendant did call 120 in a timely manner, but this action was not a report to the authorities. After learning that the victim had died, the defendant fled the scene. His actions should be recognized as <i>fleeing</i> , and his defense is not accepted.

Table 6: Cases mentioning the vague concept and Cases discussing in detail why the vague concept applies. We only consider the latter as the relevant cases.

C Details of Automatic evaluation

C.1 Implementation Details

We filtered a total of 2,642 cases and extracted 2,642 reasons for generating concept interpretations. On average, each concept was associated with 165 cases. We use the open-source LLM Qwen2.5-72B-Instruct with a maximum context length of 128k tokens to generate vague concept interpretations. The temperature is set to 0.9 to encourage more diverse outputs. Detailed prompt information can be found in Appendix F.1.4.

To investigate the effectiveness of our generated interpretations in assisting models with different capabilities, we employ Qwen2.5-72B-Instruct and Qwen2.5-14B-Instruct to perform the automatic evaluation task.

To reduce the randomness of the output, the temperature of all LLMs for prediction is set to 0, and the generation process is repeated three times. Among the predictions, we select the label $\hat{l}$ that appears most frequently. From the responses associated with $\hat{l}$, one is randomly chosen, and its reason $\hat{r}$ is extracted for consistency scoring. We use gpt-4o-2024-08-06 (Achiam et al., 2023) to give the consistency score, setting the temperature to 0.

C.2 Manual inspection of the LLM-annotated data

To evaluate the relevance between the LLM filtered case judgments and the vague concepts, we randomly sampled 20 cases for each concept from $\mathcal{D}$ and manually assessed their relevance to the vague concepts. The results show that over 96% of the cases are indeed relevant to the vague concepts. In addition, manual inspection of 200 extraction results indicates that the accuracy of Qwen2.5-72B-Instruct in labeling the gold label l and the reasoning r are 98% and 94%, respectively.

C.3 Example of gold labels and gold reasons

Table 7 shows some examples of gold labels and gold reasons.

Label	Reason
Yes	The location of the theft is a closed store that integrates living quarters and business operations. Since the store is connected to the living area, and after closing, it becomes part of the living space, relatively isolated from the outside, this theft is classified as theft by entering a dwelling.
No	The dormitory is a collective dormitory of the factory, intended solely for employees to rest during lunch breaks and nighttime. It does not include facilities for dining or other living functions and lacks the characteristics of a dwelling. Therefore, the accusation of the defendant committing theft by entering a dwelling is inappropriate.

Table 7: Examples of gold labels and their corresponding gold reasons .

C.4 Detailed Results

C.4.1 Different Models

As shown in Table 8, to validate the generalizability of our method, we utilized different LLMs to generate interpretations and perform automatic evaluations. Due to the cost constraints of APIs, we conducted experiments on a subset of our test dataset. Our findings are as follows: (1) **Stronger models demonstrate greater ability to generate concept interpretations** . The interpretations generated using Qwen2.5 (72B) and GPT-4o lead to noticeably higher performance improvements than using GPT-4o-mini. (2) **Generated concept interpretations can assist even weaker LLMs in accurately understanding vague concepts**. In our method, the performance gap between GLM and the other models is significantly smaller than that observed in the Zero-shot baseline.

Interpret model	Qwen2.5 (72B)			gpt-4o-2024-08-06			gpt-4o-mini		
Predict model	Qwen	GPT	GLM	Qwen	GPT	GLM	Qwen	GPT	GLM
Zero-Shot	57.27	51.68	47.06	57.27	51.68	47.06	57.27	51.68	47.06
Direct Interpretation	61.58	53.65	53.14	61.02	52.70	54.96	55.94	51.80	50.15
Judicial Interpretation	62.14	59.05	53.05	62.14	59.05	53.05	62.14	59.05	53.05
ATRI (Ours)	66.67	59.01	60.34	61.99	60.01	59.23	63.14	54.14	54.18

Table 8: Macro-F1 results of using different LLMs to generate interpretations and perform the Legal Concept Entailment Task on a subset. Here, **Qwen**, **GPT**, and **GLM** represent Qwen2.5-72B-Instruct, gpt-4o-mini, and GLM-4-9B-Chat(GLM et al., 2024), respectively.

C.4.2 Model Bias

Analyzing the LLM’s predictions reveals a strong bias toward responding with "Yes" on the Legal Concept Entailment Task.

D Details about human evaluation

D.1 Details about evaluation metrics

- **Accuracy (Acc.)** The interpretation should align with the current legal articles and relevant judicial interpretations, avoiding any misinterpretation or distortion of the original intent of the law.
- **Informativeness (Info.)** The interpretation should provide additional insights that were previously unknown, thereby enhancing the human evaluators’ legal knowledge beyond their prior understanding.

	Qwen2.5 (72B)			Qwen2.5 (14B)		
	Pos	Neg	Ratio	Pos	Neg	Ratio
Zero-Shot	2285	367	6.23	2329	323	7.21
Chain-of Thought	2216	436	5.08	2313	338	6.84
Direct Interpretation	1989	662	3.00	2049	602	3.40
Judicial Interpretation	2018	634	3.18	2011	641	3.14
ATRI	1939	713	2.72	1926	726	2.65
Gold label	1714	837	2.05	1714	837	2.05

Table 9: The number and ratio of positive and negative cases predicted by the LLM. *Pos* represents the number of cases predicted as "Yes", *Neg* represents the number of cases predicted as "No", and *Ratio* denotes the ratio of *Pos* to *Neg*.

- **Normativity (Norm.)** The interpretation should conform to the standard expressions and terminology used within the field of legal studies.
- **Comprehensiveness (Comp.)** The interpretation should cover as many relevant scenarios as possible, including applicable and excluded cases, ensuring that no key aspects are omitted.
- **Readability (Read.)** The interpretation should be expressed in clear, simple language, avoiding excessive use of legal jargon or complex sentence structures, so that even non-experts can generally understand the meaning and application of the legal concept.

E Details of the generated concept interpretation

E.1 The structure of generated concept interpretation

The generated concept interpretation includes the following main components:

- **Analysis:** Cites judicial interpretations to define the basic meaning, applicability conditions, and exclusions of the vague concept.
- **Example Cases:** Provides specific case examples illustrating how the vague concept is applied; this section includes 5 Positive Cases and 5 Negative Cases.
- **Judicial Discretion:** Provides multiple judgment criteria to guide judges on how to flexibly apply the vague concept based on the specifics of the case.

E.2 Details of the interpretation example e_0

We additionally select a vague concept c_0 and its corresponding article a_0 . c_0 and a_0 are not the same as any of the concepts and articles selected in Section ???. Using the methods outlined in previous sections, we derive a reason set $\mathcal{R}_0$. These three components serve as input to the LLM. We generate multiple distinct interpretations. A legal expert then selects one interpretation that best adheres to legal format specifications and modifies it to ensure correctness and clarity. We designate the revised interpretation as the interpretation example e_0 .

E.3 An example of generated vague concept interpretation

E.3.1 Original text in Chinese

在中华人民共和国刑法第二百六十四条中，“盗窃公私财物，数额较大的”涉及盗窃行为的定罪和量刑，该条文的实施中，其中的“入户盗窃”中“户”的概念可能会产生一定的法律解释上的模糊性。司法程序中，法官需要根据案件的实际情况对“户”的定义进行具体化和解释。

解析

1. **基本定义**:

- 根据最高人民法院、最高人民检察院《关于办理盗窃刑事案件适用法律若干问题的解释》,“户”的特征表现为供他人家庭生活和与外界相对隔离的两个方面。
- “户”通常包括家庭的居住场所、封闭的院落、为生活租用的房屋等。
- 非法进入他人生活区域与外界相对隔离的住所盗窃的,应当认定为“入户盗窃”。

2. **具体适用**:

- 对于“户”进行具体适用时,需要查看被盗场所是否符合供他人家庭生活的场所,并且与外界相对隔离。
- 对于公共场所、商业用途的场所或者未经明确隔离的区域,一般不被认定为“户”。
- 在具体案件中,法官会根据房屋的用途、侵入方式、时间等切实情况进行判断。

3. **排除情况**:

- 不符合“生活用途”:如仅为商业用途的店铺、公共办公场所等。
- 不具备“相对隔离性”:如无任何封闭、开放性极强的场所。
- 他人同意或者空置:如经允许进入的情况下进行盗窃,或者在实际无人生活的装修或空房中进行盗窃。

举例说明

- **符合“户”定义的案例**:

1. **案例一**:

- **具体情形**:被告人非法进入供他人家庭生活的封闭住所进行盗窃行为。
- **案例说明**:被告人余某甲非法进入xx区xxx村xxx号305室,该305室是他人租住的住宅,具有供家庭生活和与外界相对隔离的特征,符合“户”的定义。
- **判决结果**:法院认定其为入户盗窃,因其非法进入相对隔离的私人住宅内实施盗窃。

2. **案例二**:

- **具体情形**:被告人多次进入他人家庭住所,在家人不在场的情况下进行盗窃。
- **案例说明**:被告人李某某的两次盗窃行为发生在被害人的住宅内,该住宅具有供家庭生活和与外界相对隔离的特征,符合“户”的定义。
- **判决结果**:法院认定其为入户盗窃,因其非法进入供家庭生活的住所。

3. **案例三**:

- **具体情形**:被告人深夜翻墙进入与外界隔离的家庭院落,并进入室内实施盗窃。
- **案例说明**:被告人田某深夜侵入多户被害人家中实施盗窃,这些住所均符合供家庭生活和与外界相对隔离的特征。
- **判决结果**:法院认定其为入户盗窃,因其非法进入家庭生活用的封闭场所。

4. **案例四**:

- **具体情形**:被告人利用工具撬锁,破门进入封闭的私人住所实施盗窃。
- **案例说明**:被告人张某某利用窃取的钥匙进入被害人黄某某家中实施盗窃,该住宅具有供家庭生活和与外界相对隔离的特征。
- **判决结果**:法院认定其为入户盗窃,因其非法进入私人家庭住所。

5. **案例五**:

- **具体情形**:被告人在家人经常出入的生活区域安静时间段入内盗窃。
- **案例说明**:被告人王某某多次采用秘密手段窃取公民财物,且其行为发生在户内,即被害人的住宅内。
- **判决结果**:法院认定其为入户盗窃,因其非法进入供他人家庭生活且与外界相对隔离的场所。

- **不符合“户”定义的案例**:	899
1. **案例一**:	900
- **具体情形**:	901
被告人盗窃商业用途的未居住店铺内的财物或者在公共区域内实施盗窃。	902
- **案例说明**:	903
被告人刘某某在被害人经营的商铺实施盗窃,而非进入被害人家庭生活的住所。	904
- **判决结果**:	905
法院认定其不属于入户盗窃,因为商铺主要用于商业经营,不符合“户”的定义。	906
2. **案例二**:	907
- **具体情形**:	908
被告人在装修未居住的房屋中实施盗窃行为。	909
- **案例说明**:	910
被告人张某某盗窃的场所是出租楼一楼用于停放车辆的公共场所,不属于严格意义上的户。	911
- **判决结果**:	912
法院认定其不属于入户盗窃,因为该房屋未用于居住且不符合“户”的定义。	913
3. **案例三**:	914
- **具体情形**:	915
被告人在被害人用来经营的场所内盗窃,但该场所主要功能为商业用途并不具备生活属性。	916
- **案例说明**:	917
被告人刘某某进入悬挂“中国移动通信/雨露/指定专营店”的店铺,该店位于被害人家庭所有房屋的一层,一层前部区域为手机经营区,一层后部及楼上区域为曾某家居生活区,案发时一楼营业区与生活区被墙、门明确隔离。	918
- **判决结果**:	919
法院认定其不属于入户盗窃,因为被盗场所主要用于商业经营,且与生活区明确隔离。	920
4. **案例四**:	921
- **具体情形**:	922
被告人公共办公用途的建筑物进行盗窃,但未达到与外界相对隔离和家庭生活双重特征的场所。	923
- **案例说明**:	924
被告人王某某盗窃进入的场所是被害人所在单位的职工宿舍,主要用于职工工作之余的休息,不具备家庭生活功能。	925
- **判决结果**:	926
法院认定其不属于入户盗窃,因为职工宿舍主要用于休息,不具备家庭生活属性。	927
5. **案例五**:	928
- **具体情形**:	929
被害人不居住的长期搁置用房,被犯罪嫌疑人在这种场所中的盗窃行为。	930
- **案例说明**:	931
被告人朱某甲盗窃的场所是xx小区的一处房屋,该房屋自2012年起无人居住,不符合供家庭生活使用的特征。	932
- **判决结果**:	933
法院认定其不属于入户盗窃,因为该房屋未用于任何生活活动,不符合“户”的定义。	934
### 司法裁量	935
法官在具体案件中对“户”的定义的司法裁量主要基于以下几点:	936
- **实际用途**:	937
确认被盗场所是否用于家庭生活。	938
- **场所性质**:	939
确认场所有没有明确的与外界隔离的措施,如围墙、门等。	940
- **生活特征**:	941
检查场所是否具备日常家庭生活的基本设施。	942
- **入侵方式**:	943
查看被告人是否非法进入有明确隔离的家庭生活场所。	944
- **时间特征**:	945
如家庭成员日常活动时间被侵入更容易被认定为“户”。	946
	947
	948
	949
	950

951 综上，“户”在独立案件中需结合案件场所的用途、隔离性和实际生活特征进行认定，司法审查
952 需要严密把握其适用性以确保判决的法律公正。

953 E.3.2 English translation

954 Article 264 of the Criminal Law of the People's Republic of China states: "Theft of
955 public or private property in large amounts" pertains to the conviction and sentencing
956 of theft offenses. The implementation of this article involves a potential ambiguity
957 in the legal interpretation of the term "dwelling" in the context of "burglary into
958 a dwelling." In judicial procedures, judges need to concretely define and interpret
959 the concept of "dwelling" based on the specific facts of each case.

960 ### Analysis

961 1. **Basic Definition**:

962 - According to the "Interpretation on the Application of Law in Handling Theft
963 Criminal Cases" by the Supreme People's Court and the Supreme People's Procuratorate,
964 the concept of "dwelling" is characterized by two aspects: providing for family life
965 and being relatively isolated from the outside world.

966 - "Dwelling" typically includes residential premises, enclosed courtyards, rented
967 homes used for living, etc.

968 - Theft involving illegal entry into someone else's residence, which is isolated
969 from the outside world, should be considered "burglary into a dwelling."

970 2. **Specific Application**:

971 - When applying the concept of "dwelling," it is necessary to check whether the
972 stolen premises are used for family life and are isolated from the outside.

973 - Public places, commercial premises, or areas without clear isolation are
974 generally not considered "dwellings."

975 - In specific cases, the judge will assess factors such as the purpose of the
976 property, the manner of intrusion, and the time of the offense.

977 3. **Exclusions**:

978 - Not fitting the "living purpose": Such as shops used solely for business, public
979 offices, etc.

980 - Lack of "relative isolation": For example, places with no enclosure or highly
981 open areas.

982 - Consent or vacancy: If theft occurred with permission or in vacant, uninhabited
983 buildings under renovation.

984 ### Example Cases

985 - **Cases that meet the definition of "dwelling"**:

986 1. **Case 1**:

987 - **Facts**: The defendant illegally entered a private residence used for family
988 living.

989 - **Explanation**: The defendant, Mr. Yu, unlawfully entered Room 305 of Building
990 XXX in Village XXX, District XX, which is rented by another person and used for family
991 life, isolated from the outside world. This meets the definition of "dwelling."

992 - **Verdict**: The court ruled it as burglary into a dwelling, as the defendant
993 unlawfully entered a private residence that was relatively isolated.

2. Case 2 :	1002
- Facts : The defendant entered a family home repeatedly while the residents were absent.	1003
- Explanation : The defendant, Mr. Li, committed two thefts in the residence of the victim, which was used for family life and isolated from the outside. This meets the definition of "dwelling."	1004
- Verdict : The court ruled it as burglary into a dwelling because the defendant illegally entered a residential property used for family living.	1005
	1006
	1007
	1008
	1009
	1010
3. Case 3 :	1011
- Facts : The defendant climbed over a wall to enter a family courtyard isolated from the outside world and then committed theft.	1012
- Explanation : The defendant, Mr. Tian, illegally entered several victims' homes late at night. These homes were used for family life and were isolated from the outside world.	1013
- Verdict : The court ruled it as burglary into a dwelling because the defendant unlawfully entered a family living space that was enclosed.	1014
	1015
	1016
	1017
	1018
	1019
4. Case 4 :	1020
- Facts : The defendant used tools to pry open a lock and break into a private residence to commit theft.	1021
- Explanation : The defendant, Mr. Zhang, used stolen keys to enter the victim's home to commit theft. This residence was used for family life and isolated from the outside.	1022
- Verdict : The court ruled it as burglary into a dwelling because the defendant unlawfully entered a private home.	1023
	1024
	1025
	1026
	1027
	1028
5. Case 5 :	1029
- Facts : The defendant entered a residential area during a time when family members frequently came and went.	1030
- Explanation : The defendant, Mr. Wang, repeatedly stole property from a family residence using secretive methods. His actions occurred inside the victim's home, which was a residential space.	1031
- Verdict : The court ruled it as burglary into a dwelling because the defendant illegally entered a residential area used for family life and isolated from the outside.	1032
	1033
	1034
	1035
	1036
	1037
	1038
- Cases that do not meet the definition of "dwelling" :	1039
	1040
1. Case 1 :	1041
- Facts : The defendant stole property from a commercial store or in a public area.	1042
- Explanation : The defendant, Mr. Liu, committed theft in a shop operated by the victim, which was not a family residence.	1043
- Verdict : The court ruled it was not burglary into a dwelling because the shop was primarily for commercial use, not for family living.	1044
	1045
	1046
	1047
	1048
2. Case 2 :	1049
- Facts : The defendant committed theft in an uninhabited property under renovation.	1050
- Explanation : The defendant, Mr. Zhang, stole from a public space used for vehicle parking in a building that was not a residential area.	1051
	1052
	1053

- **Verdict**: The court ruled it was not burglary into a dwelling because the property was not used for living purposes.

3. **Case 3**:

- **Facts**: The defendant committed theft in a commercial space that did not serve residential purposes.

- **Explanation**: The defendant, Mr. Liu, entered a shop (labeled "China Mobile/ Yue Lu/ Designated Specialty Store") on the first floor of a building owned by the victim. The front area of the first floor was a commercial section selling mobile phones, while the rear and upper floors were residential areas. At the time of the offense, the commercial and residential areas were clearly separated by walls and doors.

- **Verdict**: The court ruled it was not burglary into a dwelling because the stolen property was in a commercial space, separate from the residential area.

4. **Case 4**:

- **Facts**: The defendant entered a public office building to commit theft, but the location did not have the characteristics of a dwelling.

- **Explanation**: The defendant, Mr. Wang, entered the dormitory of the victim's workplace, which was used by employees for rest, not for family living.

- **Verdict**: The court ruled it was not burglary into a dwelling because the dormitory was used for rest and not for family living.

5. **Case 5**:

- **Facts**: The defendant stole from a long-term uninhabited property.

- **Explanation**: The defendant, Mr. Zhu, committed theft in a house in the XX community that had been uninhabited since 2012 and was not used for family living.

- **Verdict**: The court ruled it was not burglary into a dwelling because the property was not used for living activities and did not meet the definition of "dwelling."

Judicial Discretion

Judges' judicial discretion in defining "dwelling" in specific cases mainly relies on the following factors:

- **Actual Use**: Confirming whether the stolen property was used for family life.

- **Nature of the Residence**: Confirming whether the residence had clear isolation measures such as walls or doors.

- **Living Features**: Checking whether the premises had basic facilities for daily family life.

- **Intrusion Method**: Determining whether the defendant illegally entered a clearly isolated family living space.

- **Time Features**: For instance, when family members' daily activities are disrupted, it is more likely to be recognized as a "dwelling."

In conclusion, the definition of "dwelling" in individual cases needs to be based on the use, isolation, and actual living characteristics of the premises. Judicial review requires careful attention to ensure the proper legal application and fairness of the verdict.

F Prompts	1103
F.1 Original text in Chinese	1104
F.1.1 Prompt for determining whether court view provides a specific reason	1105
法律语言具有模糊性，而司法程序是对立法语言的一个明晰过程。在部分案件中，法官会根据案件事实对法律条文中的模糊概念进行具体化并在裁判文书中的“法庭观点”部分给出认定理由。我们考虑法条“ <code>{{article}}</code> ”中的模糊概念“ <code>{{concept}}</code> ”。我将给你一段法庭观点，请你判断法庭观点中，是否存在具体的句子解释“ <code>{{concept}}</code> ”适用或不适用于该案件的原因。先输出你的判断理由，然后严格按照以下格式输出你的最终判断。如果法庭观点中存在解释“ <code>{{concept}}</code> ”是否适用的句子，输出“ <code>[[是]]</code> ”；否则，输出“ <code>[[否]]</code> ”。	1106 1107 1108 1109 1110 1111
[法庭观点]	1112
<code>{{court view}}</code>	1113
F.1.2 Prompt for classifying whether concept c applies or not	1114
法律语言具有模糊性，而司法程序是对立法语言的一个明晰过程，法官会根据案件事实对法律条文中的模糊概念进行具体化并在裁判文书中的“法庭观点”部分给出认定理由。我们考虑法条“ <code>{{article}}</code> ”中的模糊概念“ <code>{{concept}}</code> ”。我将给你一段裁判文书中的法庭观点，请你判断法官认为模糊概念“ <code>{{concept}}</code> ”是否适用于案件中的情况。先给出你的判断理由，然后严格按照以下格式输出你的最终判断：如果“ <code>{{concept}}</code> ”适用于案件中的情况，输出“ <code>[[是]]</code> ”；否则，输出“ <code>[[否]]</code> ”。	1115 1116 1117 1118 1119 1120
[法庭观点]	1121
<code>{{court view}}</code>	1122
F.1.3 Prompt for extracting reason r from court view	1123
法律语言具有模糊性，而司法程序是对立法语言的一个明晰过程。法官会根据案件事实对法律条文中的模糊词进行具体化并在裁判文书中的“法庭观点”部分进行分析。在法条“ <code>{{article}}</code> ”中，模糊概念是“ <code>{{concept}}</code> ”。请你阅读裁判文书中的法庭观点，提取出法官对模糊概念的认定理由。理由包括对案件事实经过的分析和最后的结论。比如，如果模糊概念是“户”，你需要提取出法官认为案件中的场所满足或不满足“户”的理由是什么。	1124 1125 1126 1127 1128
[法庭观点]	1129
<code>{{court view}}</code>	1130
F.1.4 Prompt for generating concept interpretation	1131
法律语言具有模糊性，而司法程序是对立法语言的一个明晰过程。法官会根据案件事实对法律条文中的模糊概念进行具体化并在裁判文书中分析模糊概念是否适用。请你阅读给出的JSON数据，对法条中的模糊概念进行解释。其中，“法条”是待分析的模糊概念所属的法条。“模糊概念”是你需要生成解释的法律概念。“参考文本”是从许多裁判文书中提取出的解释模糊概念的文	1132 1133 1134 1135 1136
本。	1137
{	1138
"法条": <code>{{article}}</code> ,	1139
"模糊概念": <code>{{concept}}</code>	1140
"参考文本": <code>{{reasons}}</code>	1141
}	1142
以下是一个概念解释的样例，请以相同的格式规范输出。	1143
<code>{{Interpretation Example}}</code>	1144
F.1.5 Prompt for assigning consistency scores	1145
请你参考法庭观点中对“ <code>{{crime}}</code> ”中的模糊概念“ <code>{{concept}}</code> ”的认定理由，对下面模型生成的认定理由的一致性进行1-10的打分。1分代表模型生成的认定理由和法庭观点中理由完全不一	1146

致，10分代表模型生成的认定理由和法庭观点中理由完全一致。请你先输出打分理由，然后以下列格式输出你的分数：[[n]]，其中n为你的分数。

[模型生成的理由]
{{generated reason}}

[法庭观点中理由]
{{gold reason}}

F.1.6 Prompt for completing Legal Concept Entailment task

法律语言具有模糊性，而司法程序是对立法语言的一个明晰过程。法官会根据案件事实对法律条文中的模糊概念进行具体化并在裁判文书中的“法庭观点”部分分析模糊概念是否适用。在法条“{{article}}”中，模糊概念是“{{concept}}”。请你阅读下面对模糊概念的解释，根据裁判文书中的事实描述，判断案件中的情况是否适用于模糊概念“{{concept}}”。先提供判定理由，然后严格按照以下格式输出你的最终判断：如果符合模糊概念“{{concept}}”的定义，输出“[[是]]”，否则输出“[[否]]”。

[模糊概念的解释]
{{interpretation}}

[事实描述]
{{fact}}

F.2 English translation

F.2.1 Prompt for determining whether court view provides a specific reason

Legal language is inherently vague, and the judicial process serves as a clarification of legislative language. In some cases, judges may concretize vague terms in the legal texts based on the facts of the case and provide reasons for their determination in the "court view" section of the ruling document. We consider the vague concept "{{concept}}" in the legal article "{{article}}". I will give you a segment of the court view; please determine whether there is a specific sentence in the court view that explains the reason why "{{concept}}" does or does not apply to the case. First, output your reasoning for the judgment, then strictly follow the format below for your final conclusion. If there is a sentence explaining whether "{{concept}}" applies, output "[[Yes]]"; otherwise, output "[[No]]".

[Court View]
{{court view}}

F.2.2 Prompt for classifying whether concept c applies or not

Legal language is inherently vague, and the judicial process serves as a clarification of legislative language, where judges can concretize vague terms in legal texts based on the facts of the case and provide reasons for their determination in the "court view" section of the ruling document. We consider the vague concept "{{concept}}" in the legal article "{{article}}". I will give you a segment of the court view; please determine whether the judge believes the vague concept "{{concept}}" applies to the situation in the case. First, provide your reasoning for the judgment, then strictly follow the format below for your final conclusion: If "{{concept}}" applies to the situation in the case, output "[[Yes]]"; otherwise, output "[[No]]".

[Court View]
{{court view}}

F.2.3 Prompt for extracting reason r from court view

Legal language is inherently vague, and the judicial process serves as a clarification of legislative language. Judges can concretize vague terms in legal texts based on the facts of the case and analyze them in the

"court view" section of the ruling document. In the legal article "{{article}}", the vague concept is "{{concept}}". Please read the court view in the ruling document and extract the judge's reasoning for the determination of the vague concept. The reasoning includes the analysis of the facts of the case and the final conclusion. For example, if the vague concept is "dwelling," you need to extract the reasons why the judge believes the place in the case satisfies or does not satisfy the "dwelling" criterion.

[Court View] 1195
{{court view}} 1196

F.2.4 Prompt for generating concept interpretation 1197

Legal language is inherently vague, and the judicial process serves as a clarification of legislative language. Judges can concretize vague terms in legal texts based on the facts of the case and analyze whether the vague concept applies in the ruling document. Please read the given JSON data and interpret the vague concept in the legal article. Among them, "article" is the legal article to which the vague concept belongs. "vague concept" is the legal concept you need to interpret. "Reference text" is the text extracted from many ruling documents explaining the vague concept.

```
{ 1205
  "Article": {{article}}, 1206
  "vague concept": {{concept}} 1207
  "Reference text": {{reasons}} 1208
}
```

Below is an example of a concept interpretation. Please format your output following the same standard.

{{Interpretation Example}} 1212

F.2.5 Prompt for assigning consistency scores 1213

Please refer to the reasons for determining the vague concept "{{concept}}" in "{{crime}}" from the court view and rate the consistency of the following model-generated reasons on a scale of 1-10. A score of 1 indicates that the model-generated reasons are completely inconsistent with the reasons in the court view, while a score of 10 indicates complete consistency. First, output your reasoning for the score, then output your score in the following format: [[n]], where n is your score.

[Model-generated Reason] 1219
{{generated reason}} 1220

[Reason in Court View] 1221
{{gold reason}} 1222

F.2.6 Prompt for completing Legal Concept Entailment task 1223

Legal language is inherently vague, and the judicial process serves as a clarification of legislative language. Judges can concretize vague terms in legal texts based on the facts of the case and analyze them in the "court view" section of the ruling document to determine whether the vague concept applies. In the legal article "{{article}}", the vague concept is "{{concept}}". Please read the following explanation of the vague concept, and based on the factual description in the ruling document, determine whether the situation in the case applies to the vague concept "{{concept}}". First, provide reasons for your determination, then strictly follow the format below for your final conclusion: If it meets the definition of the vague concept "{{concept}}", output "[[Yes]]"; otherwise, output "[[No]]".

[Explanation of vague Concept] 1232
{{interpretation}} 1233

[Factual Description] 1234
{{fact}} 1235

1236
1237
1238

G Details of vague concepts

Tables 11 and 12 present the vague concepts we interpret and their corresponding legal articles. Table 10 presents the detailed statistics of the test set for the legal concept entailment task.

Test Dataset	
# Concepts	16
# Cases	2652
- positive	1714
- negative	837
# Average court view length	653.1
# Average fact length	4787.9
# Average reason length	160.5

Table 10: Basic statistics of the test dataset.

Vague concept	Article
情节严重	第一百二十五条：非法制造、买卖、运输、邮寄、储存枪支、弹药、爆炸物的，处三年以上十年以下有期徒刑；情节严重的，处十年以上有期徒刑、无期徒刑或者死刑。非法制造、买卖、运输、储存毒害性、放射性、传染病病原体等物质，危害公共安全的，依照前款的规定处罚。单位犯前两款罪的，对单位判处罚金，并对其直接负责的主管人员和其他直接责任人员，依照第一款的规定处罚。
情节严重	第一百二十八条：违反枪支管理规定，非法持有、私藏枪支、弹药的，处三年以下有期徒刑、拘役或者管制；情节严重的，处三年以上七年以下有期徒刑。依法配备公务用枪的人员，非法出租、出借枪支的，依照前款的规定处罚。依法配置枪支的人员，非法出租、出借枪支，造成严重后果的，依照第一款的规定处罚。单位犯第二款、第三款罪的，对单位判处罚金，并对其直接负责的主管人员和其他直接责任人员，依照第一款的规定处罚。
逃逸	第一百三十三条：违反交通运输管理法规，因而发生重大事故，致人重伤、死亡或者使公私财产遭受重大损失的，处三年以下有期徒刑或者拘役；交通运输肇事后逃逸或者有其他特别恶劣情节的，处三年以上七年以下有期徒刑；因逃逸致人死亡的，处七年以上有期徒刑。在道路上驾驶机动车，有下列情形之一的，处拘役，并处罚金：（一）追逐竞驶，情节恶劣的；（二）醉酒驾驶机动车的；（三）从事校车业务或者旅客运输，严重超过额定乘员载客，或者严重超过规定时速行驶的；（四）违反危险化学品安全管理规定运输危险化学品，危及公共安全的。机动车所有人、管理人对前款第三项、第四项行为负有直接责任的，依照前款的规定处罚。有前两款行为，同时构成其他犯罪的，依照处罚较重的规定定罪处罚。第一百三十三条之二对行驶中的公共交通工具的驾驶人员使用暴力或者抢控驾驶操纵装置，干扰公共交通工具正常行驶，危及公共安全的，处一年以下有期徒刑、拘役或者管制，并处或者单处罚金。前款规定的驾驶人员在行驶的公共交通工具上擅离职守，与他人互殴或者殴打他人，危及公共安全的，依照前款的规定处罚。有前两款行为，同时构成其他犯罪的，依照处罚较重的规定定罪处罚。
严重情节	第二百二十四条：有下列情形之一，以非法占有为目的，在签订、履行合同过程中，骗取对方当事人财物，数额较大的，处三年以下有期徒刑或者拘役，并处或者单处罚金；数额巨大或者有其他严重情节的，处三年以上十年以下有期徒刑，并处罚金；数额特别巨大或者有其他特别严重情节的，处十年以上有期徒刑或者无期徒刑，并处罚金或者没收财产：（一）以虚构的单位或者冒用他人名义签订合同的；（二）以伪造、变造、作废的票据或者其他虚假的产权证明作担保的；（三）没有实际履行能力，以先履行小额合同或者部分履行合同的方法，诱骗对方当事人继续签订和履行合同的；（四）收受对方当事人给付的货物、货款、预付款或者担保财产后逃匿的；（五）以其他方法骗取对方当事人财物的。组织、领导以推销商品、提供服务等经营活动为名，要求参加者以缴纳费用或者购买商品、服务等方式获得加入资格，并按照一定顺序组成层级，直接或者间接以发展人员的数量作为计酬或者返利依据，引诱、胁迫参加者继续发展他人参加，骗取财物，扰乱经济社会秩序的传销活动的，处五年以下有期徒刑或者拘役，并处罚金；情节严重的，处五年以上有期徒刑，并处罚金。
合同	第二百二十四条：有下列情形之一，以非法占有为目的，在签订、履行合同过程中，骗取对方当事人财物，数额较大的，处三年以下有期徒刑或者拘役，并处或者单处罚金；数额巨大或者有其他严重情节的，处三年以上十年以下有期徒刑，并处罚金；数额特别巨大或者有其他特别严重情节的，处十年以上有期徒刑或者无期徒刑，并处罚金或者没收财产：（一）以虚构的单位或者冒用他人名义签订合同的；（二）以伪造、变造、作废的票据或者其他虚假的产权证明作担保的；（三）没有实际履行能力，以先履行小额合同或者部分履行合同的方法，诱骗对方当事人继续签订和履行合同的；（四）收受对方当事人给付的货物、货款、预付款或者担保财产后逃匿的；（五）以其他方法骗取对方当事人财物的。组织、领导以推销商品、提供服务等经营活动为名，要求参加者以缴纳费用或者购买商品、服务等方式获得加入资格，并按照一定顺序组成层级，直接或者间接以发展人员的数量作为计酬或者返利依据，引诱、胁迫参加者继续发展他人参加，骗取财物，扰乱经济社会秩序的传销活动的，处五年以下有期徒刑或者拘役，并处罚金；情节严重的，处五年以上有期徒刑，并处罚金。
非法占有为目的	第二百二十四条：有下列情形之一，以非法占有为目的，在签订、履行合同过程中，骗取对方当事人财物，数额较大的，处三年以下有期徒刑或者拘役，并处或者单处罚金；数额巨大或者有其他严重情节的，处三年以上十年以下有期徒刑，并处罚金；数额特别巨大或者有其他特别严重情节的，处十年以上有期徒刑或者无期徒刑，并处罚金或者没收财产：（一）以虚构的单位或者冒用他人名义签订合同的；（二）以伪造、变造、作废的票据或者其他虚假的产权证明作担保的；（三）没有实际履行能力，以先履行小额合同或者部分履行合同的方法，诱骗对方当事人继续签订和履行合同的；（四）收受对方当事人给付的货物、货款、预付款或者担保财产后逃匿的；（五）以其他方法骗取对方当事人财物的。组织、领导以推销商品、提供服务等经营活动为名，要求参加者以缴纳费用或者购买商品、服务等方式获得加入资格，并按照一定顺序组成层级，直接或者间接以发展人员的数量作为计酬或者返利依据，引诱、胁迫参加者继续发展他人参加，骗取财物，扰乱经济社会秩序的传销活动的，处五年以下有期徒刑或者拘役，并处罚金；情节严重的，处五年以上有期徒刑，并处罚金。

Table 11: The 16 vague concepts and their corresponding articles used in our study. (i)

Vague concept	Article
情节严重	第二百二十五条：违反国家规定，有下列非法经营行为之一，扰乱市场秩序，情节严重的，处五年以下有期徒刑或者拘役，并处或者单处违法所得一倍以上五倍以下罚金；情节特别严重的，处五年以上有期徒刑，并处违法所得一倍以上五倍以下罚金或者没收财产：（一）未经许可经营法律、行政法规规定的专营、专卖物品或者其他限制买卖的物品的；（二）买卖进出口许可证、进出口原产地证明以及其他法律、行政法规规定的经营许可证或者批准文件的；（三）未经国家有关主管部门批准非法经营证券、期货、保险业务的，或者非法从事资金支付结算业务的；（四）其他严重扰乱市场秩序的非法经营行为。
户	第二百六十四条：盗窃公私财物，数额较大的，或者多次盗窃、入户盗窃、携带凶器盗窃、扒窃的，处三年以下有期徒刑、拘役或者管制，并处或者单处罚金；数额巨大或者有其他严重情节的，处三年以上十年以下有期徒刑，并处罚金；数额特别巨大或者有其他特别严重情节的，处十年以上有期徒刑或者无期徒刑，并处罚金或者没收财产。
职务	第二百七十一条：公司、企业或者其他单位的工作人员，利用职务上的便利，将本单位财物非法占为己有，数额较大的，处三年以下有期徒刑或者拘役，并处罚金；数额巨大的，处三年以上十年以下有期徒刑，并处罚金；数额特别巨大的，处十年以上有期徒刑或者无期徒刑，并处罚金。国有公司、企业或者其他国有单位中从事公务的人员和国有公司、企业或者其他国有单位委派到非国有公司、企业以及其他单位从事公务的人员有前款行为的，依照本法第三百八十二条、第三百八十三条的规定定罪处罚。
单位	第二百七十二條：公司、企业或者其他单位的工作人员，利用职务上的便利，挪用本单位资金归个人使用或者借贷给他人，数额较大、超过三个月未还的，或者虽未超过三个月，但数额较大、进行营利活动的，或者进行非法活动的，处三年以下有期徒刑或者拘役；挪用本单位资金数额巨大的，处三年以上七年以下有期徒刑；数额特别巨大的，处七年以上有期徒刑。国有公司、企业或者其他国有单位中从事公务的人员和国有公司、企业或者其他国有单位委派到非国有公司、企业以及其他单位从事公务的人员有前款行为的，依照本法第三百八十四条的规定定罪处罚。有第一款行为，在提起公诉前将挪用的资金退还的，可以从轻或者减轻处罚。其中，犯罪较轻的，可以减轻或者免除处罚。
情节严重	第二百八十条：伪造、变造、买卖或者盗窃、抢夺、毁灭国家机关的公文、证件、印章的，处三年以下有期徒刑、拘役、管制或者剥夺政治权利，并处罚金；情节严重的，处三年以上十年以下有期徒刑，并处罚金。伪造公司、企业、事业单位、人民团体的印章的，处三年以下有期徒刑、拘役、管制或者剥夺政治权利，并处罚金。伪造、变造、买卖居民身份证、护照、社会保障卡、驾驶证等依法可以用于证明身份的证件的，处三年以下有期徒刑、拘役、管制或者剥夺政治权利，并处罚金；情节严重的，处三年以上七年以下有期徒刑，并处罚金。在依照国家规定应当提供身份证明的活动中，使用伪造、变造的或者盗用他人的居民身份证、护照、社会保障卡、驾驶证等依法可以用于证明身份的证件，情节严重的，处拘役或者管制，并处或者单处罚金。有前款行为，同时构成其他犯罪的，依照处罚较重的规定定罪处罚。第二百八十条之二盗用、冒用他人身份，顶替他人取得的高等学历教育入学资格、公务员录用资格、就业安置待遇的，处三年以下有期徒刑、拘役或者管制，并处罚金。组织、指使他人实施前款行为的，依照前款的规定从重处罚。国家工作人员有前两款行为，又构成其他犯罪的，依照数罪并罚的规定处罚。
情节严重	第三百一十二条：明知是犯罪所得及其产生的收益而予以窝藏、转移、收购、代为销售或者以其他方法掩饰、隐瞒的，处三年以下有期徒刑、拘役或者管制，并处或者单处罚金；情节严重的，处三年以上七年以下有期徒刑，并处罚金。单位犯前款罪的，对单位判处罚金，并对其直接负责的主管人员和其他直接责任人员，依照前款的规定处罚。
情节严重	第三百四十八条：非法持有鸦片一千克以上、海洛因或者甲基苯丙胺五十克以上或者其他毒品数量大的，处七年以上有期徒刑或者无期徒刑，并处罚金；非法持有鸦片二百克以上不满一千克、海洛因或者甲基苯丙胺十克以上不满五十克或者其他毒品数量较大的，处三年以下有期徒刑、拘役或者管制，并处罚金；情节严重的，处三年以上七年以下有期徒刑，并处罚金。
情节严重	第三百五十九条：引诱、容留、介绍他人卖淫的，处五年以下有期徒刑、拘役或者管制，并处罚金；情节严重的，处五年以上有期徒刑，并处罚金。引诱不满十四周岁的幼女卖淫的，处五年以上有期徒刑，并处罚金。
情节严重	第三百八十四条：国家工作人员利用职务上的便利，挪用公款归个人使用，进行非法活动的，或者挪用公款数额较大、进行营利活动的，或者挪用公款数额较大、超过三个月未还的，是挪用公款罪，处五年以下有期徒刑或者拘役；情节严重的，处五年以上有期徒刑。挪用公款数额巨大不退还的，处十年以上有期徒刑或者无期徒刑。挪用用于救灾、抢险、防汛、优抚、扶贫、移民、救济款物归个人使用的，从重处罚。
情节严重	第三百九十条：对犯行贿罪的，处五年以下有期徒刑或者拘役，并处罚金；因行贿谋取不正当利益，情节严重的，或者使国家利益遭受重大损失的，处五年以上十年以下有期徒刑，并处罚金；情节特别严重的，或者使国家利益遭受特别重大损失的，处十年以上有期徒刑或者无期徒刑，并处罚金或者没收财产。行贿人在被追诉前主动交待行贿行为的，可以从轻或者减轻处罚。其中，犯罪较轻的，对侦破重大案件起关键作用的，或者有重大立功表现的，可以减轻或者免除处罚。为谋取不正当利益，向国家工作人员的近亲属或者其他与该国家工作人员关系密切的人，或者向离职的国家工作人员或者其近亲属以及其他与其关系密切的人行贿的，处三年以下有期徒刑或者拘役，并处罚金；情节严重的，或者使国家利益遭受重大损失的，处三年以上七年以下有期徒刑，并处罚金；情节特别严重的，或者使国家利益遭受特别重大损失的，处七年以上十年以下有期徒刑，并处罚金。单位犯前款罪的，对单位判处罚金，并对其直接负责的主管人员和其他直接责任人员，处三年以下有期徒刑或者拘役，并处罚金。

Table 12: The 16 vague concepts and their corresponding articles used in our study. (ii)