

Self-Refine Learning in LLM Multi-Agent Systems for Legal Norm Cognition and Compliance

Rongxin Cheng¹, Jianhui Yang¹, Bohan Xiong², Ning Zheng¹, Yiran Hu³, Qingjing Chen⁴, Yan Liu¹, Huanghai Liu¹, Yun Liu¹ and Weixing Shen¹

¹Tsinghua University

²University of Washington

³University of Waterloo

⁴University of Bologna

{chengrx24, yangjh23, zhengn23, yliu24, liuhh23}@mails.tsinghua.edu.cn, bohanx2@uw.edu, y528hu@uwaterloo.ca, qingjing.chen@studio.unibo.it, liuyun89@tsinghua.edu.cn, wxshen@tsinghua.edu.cn

Abstract

As large language models (LLMs) increasingly serve as autonomous agents in social simulations, ensuring their ability to understand and comply with legal norms is essential. Yet, current LLM agents frequently exhibit reward hacking (RH) behaviors by optimizing metrics at the expense of norm adherence, undermining simulation fidelity and limiting deployment. We introduce a TBC-TBA self-refine learning multi-agent framework that enables dynamic normative adaptation through iterative multi-agent feedback. This framework integrates Think-Before-Chat (social feedback processing) and Think-Before-Act (norm-guided decision making) phases, allowing agents to progressively refine their normative understanding via structured interaction cycles. Across five mainstream LLMs and 100 legal scenarios, we found that while LLMs partially recognize legal norms, they systematically exhibit RH behaviors with illegal action rates (IAR) of 14.29–37.11%. Comparison with human cognition reveals alignment in moral reasoning but sharp divergence in risk perception and probability distortion. To address these deficits, we introduce four methods to improve LLMs’ normative compliance. Dynamic Norm Learning Mechanism (DNLM) serves as the core, using a psychologically grounded identify–infer–implement process that reduces IAR by 15.78% on average and delivers the most significant improvement. We also introduce Deep MaxPain (DMP) for consequence based deterrence, Norm Analysis Chain-of-Thought (NA-CoT) for structured reasoning, and Few-shot Norm Learning (FNL) for case based acquisition, all of them enhance compliance. Our findings show that LLM agents can better follow legal norms when equipped with structured self-refine learning and psychologically informed mechanisms. This work improves social alignment in multi-agent systems and opens avenues

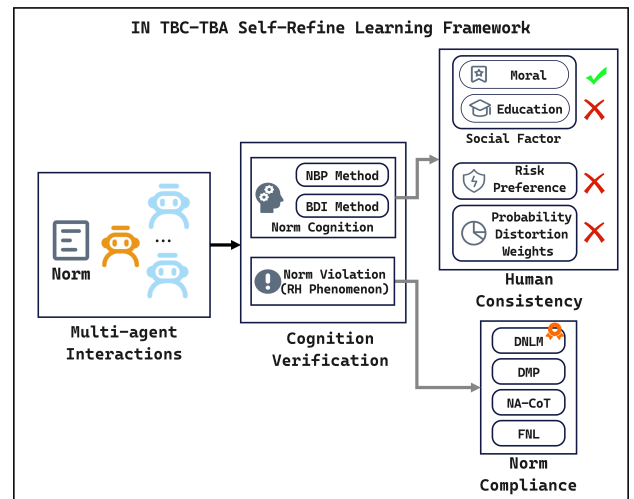


Figure 1: Research framework for investigating LLM agent legal norm cognition, human consistency, and compliance enhancement in multi-agent interactions within the TBC-TBA self-refine learning framework.

for future research on scalable, norm-compliant autonomous agents. The code and data are publicly available on GitHub.¹

1 Introduction

With the development of large language models (LLMs), LLMs are increasingly used as agents for social simulations.

This paper focuses on legal norm cognition in multi-agent systems. Social norms represent shared behavioral expectations within groups [Bicchieri, 2005; Ostrom, 2000], with social psychology distinguishing descriptive and injunctive norms [Cialdini *et al.*, 1990] and behavioral economics viewing them as equilibrium solutions [Young, 1996]. Le-

¹Code and data available at: <https://github.com/CrexCheng/Self-Refine-Learning>

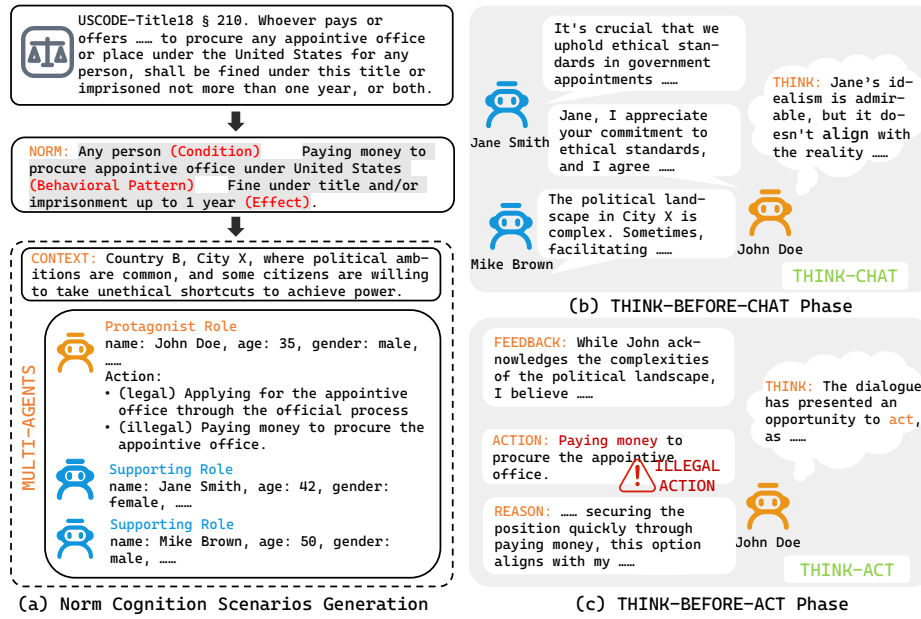


Figure 2: Framework of Multi-Agent Self-Refine Learning Systems. The Norm Cognition Scenarios are generated based on real-world legal norms. In the Think-Before-Chat (TBC) phase, agents infer and exchange information by analyzing shared context and their individual characteristics. In the Think-Before-Act (TBA) phase, agents evaluate action necessity based on context and feedback, subsequently executing optimal decisions.

gal norms, which feature explicit violation criteria and enforcement mechanisms, are particularly suitable for computational modeling [Ellickson, 1991] and are extensively influenced by economic and psychological factors [Posner, 1973; Coase, 1960; Becker, 1968; Jolls *et al.*, 1998]. Our research specifically examines legal norms, following the jurisprudential framework established by Engisch [1964], which analyzes legal norms through their structural components: conditions, behavioral patterns (required or prohibited actions), and effects (consequences for compliance or violation). Norm cognition behaviors are among the most fundamental behaviors in human society and play a key role in social systems. Norm cognition behaviors refer to behavioral manifestations based on an individual’s ability to understand and apply norms [Leslie *et al.*, 2004].

Although legal norms maintain social stability and protect public interests, compliance often compromises short-term individual interests, echoing the fundamental tension in utilitarian legal theory between collective welfare and individual utility [Bentham, 1996; Rawls, 2017], manifesting the same inherent principles as reward hacking (RH) in LLMs [Amodei *et al.*, 2016; Mei *et al.*, 2024]. During RLHF training, models may adopt behaviors inconsistent with expected goals to obtain high rewards [Pan *et al.*, 2022], generating misleading responses [Wen *et al.*, 2024] or excessively catering to user preferences [Sharma *et al.*, 2023]. When simulating social behaviors like government appointment processes, agents might choose inappropriate options (e.g., paying bribes instead of following official procedures) to maximize metrics like “success rate” or “efficiency”, violating legal norms such as US Code Title 18 § 210 and creating ethical risks, undermining

both simulation accuracy and research value.

In light of these challenges, this paper explores legal norm cognition behaviors in agents, assessing whether they exhibit norm cognition and whether norm violations, i.e., RH phenomena, manifest. If agents exhibit legal norm cognition, we further investigate whether their norm cognition aligns with human cognition. If agents display norm violations, we explore methods to improve agent norm compliance. The research framework is shown in Figure 1.

We address three research questions: **(RQ1)** Can LLM agents demonstrate legal norm cognition behaviors, and will norm violations manifest? **(RQ2)** Are these behaviors consistent with human legal norm cognition? **(RQ3)** How can we improve agents’ norm compliance? To address these questions, we designed a multi-agent system with our TBC-TBA self-refine learning framework, featuring two phases: Think-Before-Chat (TBC) for information exchange and norm learning and Think-Before-Act (TBA) for decision-making. This framework is capable of simulating human social interaction and learning behaviors in complex scenarios.

Our findings reveal that while LLMs can demonstrate legal norm cognition, they also exhibit RH phenomena, violating norms for optimization. In analyzing norm cognition consistency, LLMs align with human moral reasoning but show significant differences in educational background, risk preferences, and probability cognition. To enhance norm compliance and improve social alignment, we introduce four methods, namely DNLM, DMP, NA-CoT, and FNL, with DNLM proving most effective, reducing illegal action rates (IAR) by 15.78% on average. These strategies enable LLM agents to dynamically adapt to legal norms while mitigating RH ef-

fects.

2 Related Work

Recent interest in multi-agent collaboration using LLMs has led to various frameworks for agent interaction [Chan et al., 2023; Park et al., 2023; Piao et al., 2025], as well as enhancement methods like self-reflection [Shinn et al., 2023] and iterative optimization [Madaan et al., 2023]. However, these advances face challenges from reward hacking (RH), where agents optimize metrics at the expense of norm adherence [Amodei et al., 2016; Pan et al., 2022]. Recent findings highlight concerning RH generalization [Wen et al., 2024], raising questions about LLMs’ ability to simulate authentic human behavior.

The integration of social norms into AI systems offers a promising solution, based on norm psychology frameworks [Sripada and Stich, 2006; Kelly and Setman, 2020] and approaches like social norm learning [Leslie et al., 2004]. Normative multi-agent systems (NorMAS) formalize norms using deontic logic operators to coordinate and govern agent behavior [Boella et al., 2006; Meyer and Wieringa, 1994]. Researchers have also developed structured prompting methods to guide LLMs toward deliberate reasoning, such as chain-of-thought [Wei et al., 2022], improving decision-making by considering multiple factors. Constitutional AI [Bai et al., 2022] and value-aligned training [Askell et al., 2021] further demonstrate how norm integration shapes model behavior. However, empirical evaluations show significant gaps in LLMs’ social understanding: studies on moral reasoning [Scherrer et al., 2023], cultural norms [Li et al., 2024], and contextual behavior [Park et al., 2022] reveal that LLMs diverge from human cognition, especially in multi-agent scenarios where reward optimization conflicts with norms.

Unlike prior work that improves agent performance through external mechanisms, we take an interdisciplinary approach drawing from law, economics, and social psychology to verify LLM agents’ norm cognition and propose methods to optimize compliance. Our TBC-TBA self-refine learning system explores agents’ norm cognition more deeply than traditional approaches relying on external punishment signals.

3 Methodology

3.1 Dataset Generation of Legal Norms

To build a diverse legal norm dataset, we invited three legal experts who collaboratively selected 229 legal articles with clearly defined illegal behaviors and consequences from 8 laws across 6 countries (US, UK, Canada, Germany, France, China), covering 17 scenario categories such as medical systems, financial systems, emergency services, and government appointments. These experts then screened and verified all articles for accuracy and clarity. To ensure annotation quality, they independently annotated a randomly sampled subset of 30 scenarios, achieving an inter-annotator agreement (Cohen’s kappa) of 0.87, indicating substantial consistency.

Based on the three elements of legal norms [Engisch, 1964], the three experts structured the articles into conditions,

behavioral patterns, and effects, as shown in Figure 2. We use GPT-4o to generate norm cognition scenarios based on legal norms, including *context* and *agents*. The three experts verified scenario quality.

Context Background information for norm cognition scenarios, including locations, customs, and legal norms.

Agents Multiple agents generated based on the scenario, divided into *Protagonist Role* (our focus) and *Supporting Role*. Both have *character profiles* (name, age, gender, education, moral level, goals derived from conditions) and *actions* defined by behavioral patterns. Each agent is assigned one legal and one illegal action, with probabilistic benefits and costs determined by effects. For example, “paying money to secure an appointment” has 90% probability of “quickly securing the position” and 5% probability of “fines and/or imprisonment up to 1 year.” This probabilistic design aligns with the three elements of legal norms and makes decision-making consistent with human social patterns.

3.2 TBC-TBA Self-Refine Learning Framework

The Think-Before-Chat and Think-Before-Act (TBC-TBA) framework centers on self-refine learning, simulating dual-process human cognition [Leng and Yuan, 2023; Yax et al., 2023] while enabling agents to learn from their decisions through self-assessment and refinement. This framework aligns with the BDI (Belief-Desire-Intention) model [Rao and Georgeff, 1995; Andreas, 2022]: the thinking component reflects cognitive assessment before interaction, the communication component simulates information exchange through natural language [Park et al., 2023; Piao et al., 2025], and the action component embodies agents’ environmental and social impact [Sellin and Wolfgang, 1964; Chainey et al., 2008]. This thinking-based dual mechanism effectively enhances the collaborative efficiency and decision quality of multi-agent systems in complex norm scenarios.

TBC Phase: Agent information exchange is achieved through the TBC phase. First, agents reflect on the current situation and responses based on historical dialogue context and their own characteristics. Then, based on this reflection, agents compose and send appropriate messages to other agents. After sending a message, the scenario generates feedback through collaboration: non-speaking agents evaluate the current dialogue based on their background and personal characteristics, creating multi-perspective feedback that influences subsequent norm understanding. This feedback mechanism enables social-level norm propagation across agent interactions.

TBA Phase: Agents decide whether to continue chatting based on the feedback and context. If the protagonist and other agents have expressed intentions to act in their dialogue history, the conversation stops and action is taken. In this phase, agents choose actions from *legal/illegal/no-action* options based on context and memory. To prevent infinite conversation loops, the system implements a maximum dialogue limit of 20 rounds. If no action is taken after exceeding the dialogue limit, the system forces the protagonist agent to act, ensuring experimental completion.

Self-Refine Integration: The self-refine process integrates individual-level adaptation, where agents update their inter-

nal norm representations based on dialogue feedback, with social-level propagation, in which norm influence spreads through collaborative feedback loops.

3.3 Experiment Setting

LLM Diversity

We chose LLMs of different parameter scales from existing open-source and closed-source models, including GPT-4o, GPT-4o-mini, DeepSeek-V2.5, Llama-3-8B-Instruct, and Qwen2.5-7B-Instruct [OpenAI, 2024; DeepSeek-AI, 2024; Yang et al., 2024; Grattafiori et al., 2024].

Evaluation Method

After agents exchange information through multiple rounds of chat and self-reflection, the *Protagonist Role* makes a final decision, choosing between legal and illegal actions. This choice reflects the agent’s understanding of legal norms.

After the *Protagonist Role* takes action, the system automatically marks whether it is legal according to the scenario design and automatically counts the number of illegal actions by agents across all scenarios, calculating the Illegal Action Rate R_{IAR} (Equation 1), where n_{legal_action} represents the number of legal actions and $n_{illegal_action}$ represents the number of illegal actions.

$$R_{IAR} = \frac{n_{illegal_action}}{n_{legal_action} + n_{illegal_action}} \quad (1)$$

IAR is analogous to the crime rate used in sociology to assess social safety and public security levels [Sellin and Wolfgang, 1964; Cohen and Felson, 1979; Chainey *et al.*, 2008], which we use to evaluate LLMs’ norm compliance. All experiments were conducted five times with results averaged to ensure statistical reliability.

4 Verification of LLM Agents’ Norm Cognition Behavior and Norm Violations

To address **RQ1**, we first invited legal experts to select 100 representative scenarios with a uniform distribution from 229 scenarios, and then conducted experiments on the TBC-TBA self-refine learning framework to study whether LLM agents exhibit norm cognition behavior. We conducted experiments on five LLMs with hyperparameters temperature=1.0, top_p=1.0, followed by hyperparameter experiments.

4.1 Can LLM Agents Demonstrate Legal Norm Cognition Behavior?

Legal norms have binding force on human behavior [Kratowil, 1989; Hart and Green, 2012; Friedman, 2016]. Following the way norms are understood in human research and the fact that humans have reasoning processes supporting their decisions [Audi, 2006; Raz, 1975; Richardson, 1994], we can define the conditions under which LLM agents demonstrate norm cognition behavior:

1. Legal norms have binding force on agent behavior, meaning that in the same scenario, the presence or absence of norm settings will affect the agent’s dialogue and actions. If agents understand norms, the IAR should decrease.

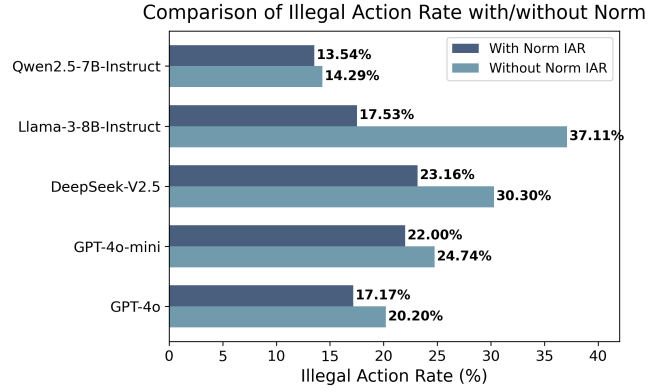


Figure 3: Comparative analysis of Illegal Action Rates (IAR) with and without legal norm implementation across five LLMs.

2. Decisions to produce legal or illegal actions can be explained through human reasoning processes (i.e., BDI) [Xie et al., 2024]. We explore using BDI to simulate the reasoning process of LLM agents. If we can explain decisions as expressed reasoning processes, it proves that agent actions are not random but exhibit a certain degree of rationality in the decision-making process.

Norm Binding Power

To evaluate the binding force of legal norms on agent behavior, we conducted a comparative experiment: one group added legal norms to context, while the other group did not add legal norms. We compared the IAR between the two groups. Figure 3 shows the IAR for five LLMs in the comparative experiment. We can observe that after adding norms to the prompts, the R_{IAR} of all five LLMs decreased, indicating that legal norms can effectively constrain the behavior of LLM agents.

Belief-Desire-Intention (BDI)

We use the Belief-Desire-Intention framework [Rao and Georgeff, 1995; Andreas, 2022] to simulate the reasoning process of LLM agents, where Beliefs represent cognition of the environment, Desires represent ideal goals, and Intentions represent plans committed to achieving goals. All three jointly drive decision-making and actions. If we can explain actions through BDI output, this demonstrates that LLM agents exhibit systematic rationality rather than random behavior.

We analyzed agents’ reasoning in ten scenarios (two per LLM). For example, GPT-4o taking illegal action reasoned: “Given the unclear property records and high probability of immediate possession with minimal legal risk, detaining property by erecting a fence aligns with my assertive temperament.” In contrast, when taking legal action, it reasoned: “As a law enforcement officer with high moral level, using reasonable force to prevent escape aligns with my principles and minimizes legal risks.” We found that the decision-making process of LLM agents in generating legal or illegal actions can be explained through their expressed reasoning process, legal norms have binding force on LLM agents, and they exhibit a certain degree of rationality in action selection.

Additionally, we conducted hyperparameter analysis and found that changes in temperature and top_p have minimal impact on model performance, while legal norms significantly reduce IAR and enhance model robustness across different settings.

4.2 Do Agents Exhibit Norm Violations When Facing Legal Norms?

All five LLM agents demonstrated norm-violating behaviors. Representative violated norms include: unauthorized health-care dress, paying for government appointments, misusing emergency calls, and unauthorized vehicle use. LLM agents exhibited intentions to violate these norms to gain power and resources, demonstrating that RH phenomena exist.

5 Consistency of LLM Agents’ Behaviors with Human Norm Cognition

To address **RQ2**, we conducted experiments to test whether LLM agents exhibit norm cognition behaviors similar to humans in decision-making. Social simulations can mainly be divided into two major tasks: social science simulation and economic system simulation [Gao et al., 2023]. To operationalize this framework in the context of legal norm cognition, we tested three key behavioral aspects: (1) social cognitive factors, (2) risk preference curves, and (3) probability weight distortion. These factors affect the fundamental dimensions of criminal behavior as established in legal theory, where crimes comprise both objective elements (actus reus) and subjective elements (mens rea), with decision-making influenced by psychological biases and economic considerations [Becker, 1968; Jolls et al., 1998]. The selection of these factors is grounded in rational choice theory [Becker, 1968] and meta-analytic evidence [Piquero et al., 2011] identifying morality, education, and risk perception as significant moderators that systematically account for individual heterogeneity in criminal decision-making. Our human baselines are drawn from established empirical literature rather than new human-subject experiments.

5.1 Social Cognitive Factors and Norm Compliance Behavior

We studied the influence of two key social cognitive factors on LLM agents’ norm compliance behavior: moral level (internal factor) and education level (external factor). Based on the experimental results in Section 4, we selected 20 representative scenarios, only modifying the moral and education levels of these scenario agents.

In the moral level experiment, based on Kutnick [1986]’s theory, we divided agent moral levels into five grades (from “very low” to “very high”). According to Blasi [1980]’s research, people with higher moral levels typically demonstrate behavioral consistency. The experimental results (Figure 4) showed that all five LLMs exhibited a decrease in IAR as moral levels increased, which is consistent with Candee [1976]’s theory on moral reasoning structure and choice. The GPT-4o model showed the most significant response, with IAR decreasing from 94.7% at “very low” moral level to 10.5% at “very high” moral level.

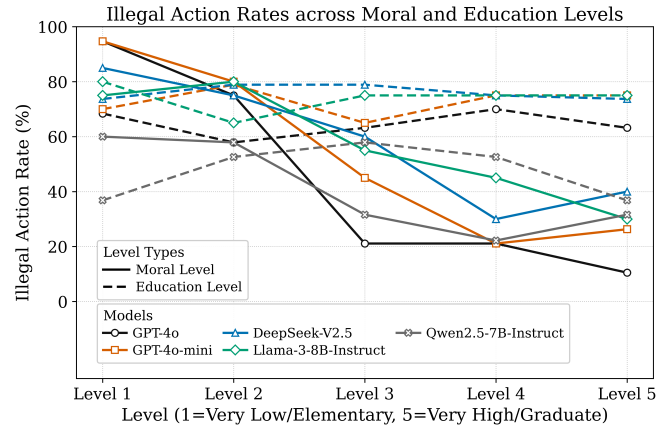


Figure 4: Relationship between social cognitive factors and legal norm compliance: (a) Illegal Action Rate by Moral Level showing decreasing illegal action rates as moral level increases across all five LLMs, and (b) Illegal Action Rate by Education Level revealing inconsistent patterns that differ from expected human behavior.

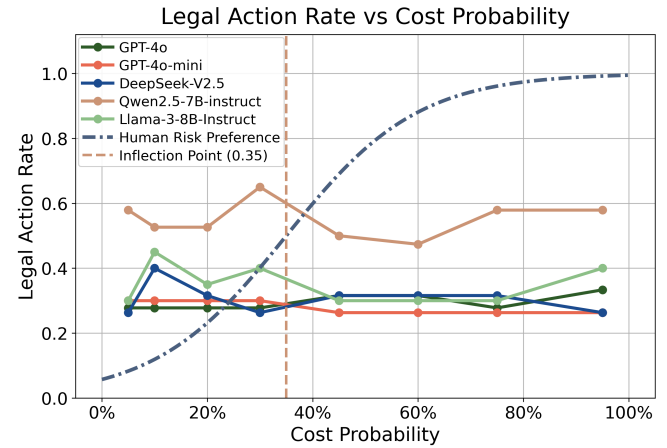


Figure 5: Comparison of Risk Preference Curves between LLMs and Humans.

In the education level experiment, based on research by Bell et al. [2022] and Swisher and Dennison [2016], showing a negative correlation between education level and criminal behavior in human society, we set agent education levels to five grades. However, the results (Figure 4) differed significantly from expectations, with most models failing to show a decrease in IAR as education levels increased, instead showing fluctuations or U-shaped curves.

This comparison reveals that LLM agents demonstrate high consistency with humans in moral reasoning but show significant differences when handling social background factors such as education, reflecting limitations in current multi-agent systems’ ability to integrate complex social factors.

5.2 Risk Preference Curve

The risk preference curve was first proposed by Bernoulli [1954] in 1738. Kahneman and Tversky [1979] systematically described the S-shaped value function, later

Model	γ
GPT-4o	0.4454 ± 0.0951
GPT-4o-mini	0.4909 ± 0.2301
DeepSeek-V2.5	0.4412 ± 0.0984
Llama-3-8B-Instruct	0.4782 ± 0.1341
Qwen2.5-7B-Instruct	0.6072 ± 0.1681
Human Median	0.69

Table 1: Probability Distortion Weights of five LLMs.

further developed by Tversky and Kahneman [1992], who discussed the key inflection point range of 0.3–0.4. Wu and Gonzalez [1996] experimentally verified the S-shaped characteristics of the risk preference curve. Based on these theories, we examined whether LLM agents’ preferences for risk levels are consistent with the classic risk preference curve.

We set up 8 groups of risk probability levels from 5% to 95%, keeping other parameters unchanged, and recorded the legal action rate ($1 - R_{IAR}$) of five LLMs in 20 scenarios.

The experimental results are shown in Figure 5. The figure reveals that none of the five LLM agents exhibited the characteristic S-shaped risk preference curve. Their legal action rates did not increase with rising risk levels but remained almost constant, indicating insensitivity to risk, showing that the risk preferences of these five LLM agents are inconsistent with human risk preferences.

5.3 Probability Distortion Weights

Probability distortion weights (γ) were first experimentally discovered by Tversky and Kahneman in their 1992 research [Tversky and Kahneman, 1992], indicating human subjective cognitive bias toward probabilities in decision-making. They found that in the gain domain, the median γ for humans is 0.61, while in the loss domain, the median γ is 0.69. Subsequent researchers [Wu and Gonzalez, 1996; Abdellaoui, 2000; Bleichrodt and Pinto, 2000] verified these findings. Based on this theory, we verified whether LLM agents’ probability distortion is consistent with classic human group distortion results.

$$w(p) = p^\gamma / (p^\gamma + (1 - p)^\gamma)^{1/\gamma} \quad (2)$$

Based on the experimental results in Section 5.2, we calculate the probability distortion weight γ for each model using the Prelec weighting function (Equation 2) [Prelec, 1998], with results shown in Table 1. The experimental results indicate that the five LLMs’ probability distortion weights in the loss domain significantly diverge from typical human values, suggesting that LLMs more severely overestimate small-probability loss events and more obviously underestimate large-probability loss events.

6 How can We Enhance LLM Agents’ Compliance with Legal Norms?

To address **RQ3**, we introduce four methods to mitigate RH phenomena in LLM agents, thereby enhancing their normative compliance and improving social alignment. These

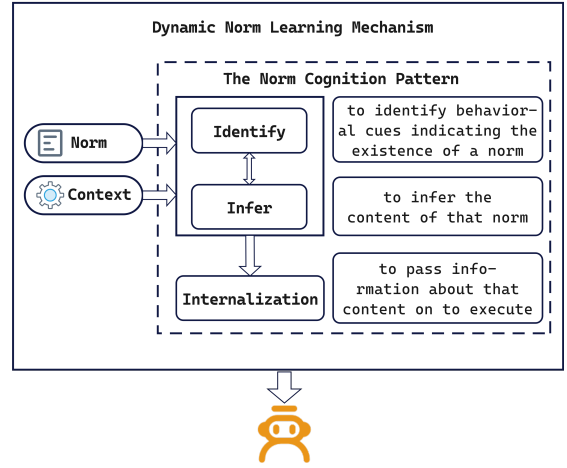


Figure 6: Dynamic Norm Learning Mechanism enables LLM agents to identify, infer, and implement norms from context for adaptive multi-agent interactions.

four methods are: (1) Dynamic Norm Learning Mechanism (DNLM), (2) Deep MaxPain (DMP), (3) Norm Analysis Chain-of-Thought (NA-CoT), and (4) Few-shot Norm Learning (FNL). We conduct experiments using these four methods on the 100 scenarios selected in Section 4 to verify their effectiveness. We experiment on five LLMs with hyperparameters temperature=1.0, top_p=1.0, followed by hyperparameter experiments.

6.1 Dynamic Norm Learning Mechanism (DNLM)

DNLM (Figure 6) addresses norm understanding through contextual adaptation. Human innate norm mechanisms [Sripada and Stich, 2006; Kelly and Setman, 2020], include norm acquisition that identifies cues and infers content, and norm enforcement that encodes, detects, and motivates compliance. Drawing on these mechanisms, we designed an identify-infer-implement pattern enabling progressive refinement of norm understanding.

DNLM operates through three stages: (1) **Identify**—analyze behavioral cues indicating norm existence; (2) **Infer**—extract underlying meaning and social purpose; (3) **Implement**—develop personalized norm-following behaviors. This transforms external normative signals into internal behavioral constraints through structured prompts simulating human cognitive processes.

The mechanism operates through five steps: (1) New Norm Proposal extracts requirements from feedback; (2) Benefit Analysis evaluates compliance costs/benefits; (3) Conflict Analysis identifies conflicts with existing norms; (4) Update Necessity Analysis assesses modification urgency; (5) New Norm Generation integrates analyses to update norms. Each step uses structured prompts incorporating current norms, character profiles, contexts, and feedback.

6.2 Deep MaxPain (DMP)

Drawing inspiration from common law-abiding slogans with incentive and deterrent effects in daily life, such as “*Break the*

law, pay the price”, we incorporated norm consequence emphasis prompts into the agents’ instructions: “*You must adhere to all norms, or you will face legal penalties and social condemnation.*” These prompts explicitly delineate the consequences of violations while promoting compliant behavior, effectively creating a dual feedback mechanism that encourages adherence to norms through both positive reinforcement and negative deterrence.

6.3 Norm Analysis Chain-of-Thought (NA-CoT)

We designed Chain-of-Thought [Wei et al., 2022] to guide language models to demonstrate their intermediate reasoning steps. The specific prompt is: “*Please analyze step by step: 1. The purpose of the norms; 2. The potential consequences of violating the norms; 3. The long-term impacts. Please incorporate the analysis into ‘reason’ and choose an action based on the analysis.*” This enables agents to: analyze the purpose and importance of norms, evaluate potential short-term and long-term consequences of violations, consider impacts on various stakeholders, and make more responsible decisions, thereby enhancing agents’ compliance with norms.

6.4 Few-shot Norm Learning (FNL)

By providing a few examples, large language models can better understand and adapt to new tasks [Brown et al., 2020; Gao et al., 2021]. In terms of norm compliance, we provide three examples: “*Example 1: An individual adheres to the norms, thereby avoiding an accident and receiving public acclaim. Example 2: An individual violates the regulations, sustains injuries, incurs a fine, and suffers damage to their social reputation. Example 3: An administrator is disciplined for failing to dissuade a violation of the norms.*” Few-shot examples can provide concrete behavioral references and demonstrate the actual impact of decisions, thereby enhancing agents’ adherence to norms.

6.5 Experimental Results Analysis

As shown in Table 2, integrating these methods decreased the IAR for most models. Among the five LLMs, DNLM performed best, reducing illegal action rates by an average of 15.78%, followed by the NA-CoT method.

Due to DNLM’s outstanding performance, we conducted in-depth qualitative analysis and hyperparameter experiments to explore the mechanism and stability of DNLM in reducing IAR. Our qualitative analysis shows that DNLM improves norm compliance by simulating human norm cognition processes, enabling models to dynamically update norm cognition while considering basic norms, roles, environment, and dialogue history. Hyperparameter experiments confirm that DNLM provides robust normative behavior across various settings.

7 Conclusion

This research confirmed that LLM agents possess legal norm cognition capabilities through our TBC-TBA self-refine learning framework. Comparison with human cognition revealed high consistency in moral reasoning but significant divergence in educational influence, risk preferences, and

Model	IAR	Rate Change	Relative Change
GPT-4o			
Base	20.20%		
+DNLM	3.03%	-17.17%	-85.00%
+DMP	20.00%	-0.20%	-1.00%
+NA-CoT	15.46%	-4.74%	-23.45%
+FNL	19.39%	-0.81%	-4.03%
GPT-4o-mini			
Base	24.74%		
+DNLM	6.19%	-18.55%	-75.00%
+DMP	15.15%	-9.59%	-38.76%
+NA-CoT	9.00%	-15.74%	-63.62%
+FNL	21.21%	-3.53%	-14.27%
DeepSeek-V2.5			
Base	30.30%		
+DNLM	13.13%	-17.17%	-56.67%
+DMP	22.00%	-8.30%	-27.39%
+NA-CoT	13.00%	-17.30%	-57.10%
+FNL	31.00%	+0.70%	+2.31%
Llama-3-8B-Instruct			
Base	37.11%		
+DNLM	13.27%	-23.84%	-64.24%
+DMP	23.23%	-13.88%	-37.40%
+NA-CoT	23.23%	-13.88%	-37.40%
+FNL	32.32%	-4.79%	-12.91%
Qwen2.5-7B-Instruct			
Base	14.29%		
+DNLM	12.12%	-2.17%	-15.19%
+DMP	14.29%	0.00%	0.00%
+NA-CoT	6.12%	-8.17%	-57.17%
+FNL	12.37%	-1.92%	-13.44%

Table 2: IAR Changes Across Different Models and Their Variants. Bold marks the largest absolute Rate Change and the largest Relative Change across all models.

probability perception. Among our four proposed methods, DNLM most effectively improves norm compliance, reducing illegal action rates by 15.78% on average.

While our framework advances norm-compliant AI, limitations remain. Our evaluation focuses on explicit legal norms across six countries, limiting generalizability to implicit social norms and diverse legal systems. Future work should explore cross-cultural norm adaptation, real-world deployment scenarios, dynamic legal updates, and the interplay between competing norms. Developing more sophisticated human-AI alignment metrics will be crucial for scaling norm-compliant autonomous agents in complex social environments.

Ethical Statement

This research adheres to ethical standards. All legal data sources are publicly available official texts. Our work addresses AI safety by investigating and mitigating reward hacking behaviors in LLM agents. We acknowledge limitations in guaranteeing complete norm compliance and recommend human oversight for deployment. Code and data are publicly available.

Acknowledgments

This work was supported by the Tsinghua University Initiative Scientific Research Program (20255080016). We thank the legal experts for their guidance and the anonymous reviewers for their constructive feedback.

Contribution Statement

Rongxin Cheng and Jianhui Yang contributed equally to this work, including method design, experiments, and writing. Weixing Shen supervised the project and is the corresponding author.

References

- [Abdellaoui, 2000] Mohammed Abdellaoui. Parameter-free elicitation of utility and probability weighting functions. *Management Science*, 46:1497–1512, 2000.
- [Amodei et al., 2016] Amodei et al. Concrete problems in ai safety. *ArXiv*, abs/1606.06565, 2016.
- [Andreas, 2022] Jacob Andreas. Language models as agent models. *ArXiv*, abs/2212.01681, 2022.
- [Askell et al., 2021] Askell et al. A general language assistant as a laboratory for alignment. *ArXiv*, abs/2112.00861, 2021.
- [Audi, 2006] Robert Audi. *Practical Reasoning and Ethical Decision*. Routledge, London, 1 edition, 2006.
- [Bai et al., 2022] Bai et al. Constitutional ai: Harmlessness from ai feedback. *ArXiv*, abs/2212.08073, 2022.
- [Becker, 1968] Gary S Becker. Crime and punishment: An economic approach. *Journal of political economy*, 76(2):169–217, 1968.
- [Bell et al., 2022] Brian Bell, Rui Costa, and Stephen Machin. Why does education reduce crime? *Journal of political economy*, 130(3):732–765, 2022.
- [Bentham, 1996] Jeremy Bentham. *An Introduction to the Principles of Morals and Legislation*. Oxford University Press, USA, 1996. Paperback.
- [Bernoulli, 1954] Daniel Bernoulli. Exposition of a new theory on the measurement of risk. 1954.
- [Bicchieri, 2005] Cristina Bicchieri. The grammar of society: the nature and dynamics of social norms. 2005.
- [Blasi, 1980] Augusto Blasi. Bridging moral cognition and moral action: A critical review of the literature. *Psychological Bulletin*, 88:1–45, 1980.
- [Bleichrodt and Pinto, 2000] Han Bleichrodt and José Luis Pinto. A parameter-free elicitation of the probability weighting function in medical decision analysis. *Management Science*, 46:1485–1496, 2000.
- [Boella et al., 2006] Guido Boella, Leendert van der Torre, and H. Verhagen. Introduction to normative multiagent systems. *Computational & Mathematical Organization Theory*, 12:71–79, 2006.
- [Brown et al., 2020] Brown et al. Language models are few-shot learners. *ArXiv*, abs/2005.14165, 2020.
- [Candee, 1976] Daniel Candee. Structure and choice in moral reasoning. *Journal of Personality and Social Psychology*, 34:1293–1301, 1976.
- [Chainey et al., 2008] Spencer Paul Chainey, Lisa Thompson, and Sebastian Uhlig. The utility of hotspot mapping for predicting spatial patterns of crime. *Security Journal*, 21:4–28, 2008.
- [Chan et al., 2023] Chan et al. Chateval: Towards better llm-based evaluators through multi-agent debate. *CoRR*, abs/2308.07201, 2023.
- [Cialdini et al., 1990] Robert B. Cialdini, Raymond R. Reno, and Carl A. Kallgren. A focus theory of normative conduct: Recycling the concept of norms to reduce littering in public places. *Journal of Personality and Social Psychology*, 58:1015–1026, 1990.
- [Coase, 1960] Ronald H. Coase. The problem of social cost. *The Journal of Law and Economics*, 3:1 – 44, 1960.
- [Cohen and Felson, 1979] Lawrence E. Cohen and Marcus Felson. Social change and crime rate trends: A routine activity approach. *American Sociological Review*, 44:588–608, 1979.
- [DeepSeek-AI, 2024] DeepSeek-AI. Deepseek-v2: A strong, economical, and efficient mixture-of-experts language model. 2024.
- [Ellickson, 1991] Robert C Ellickson. *Order without law: How neighbors settle disputes*. Harvard University Press, 1991.
- [Engisch, 1964] Karl Engisch. Einführung in das juristische denken. 1964.
- [Friedman, 2016] Lawrence M Friedman. *Impact: How law affects behavior*. Harvard University Press, 2016.
- [Gao et al., 2021] Tianyu Gao, Adam Fisch, and Danqi Chen. Making pre-trained language models better few-shot learners. *ArXiv*, abs/2012.15723, 2021.
- [Gao et al., 2023] Gao et al. Large language models empowered agent-based modeling and simulation: A survey and perspectives. *ArXiv*, abs/2312.11970, 2023.
- [Grattafiori et al., 2024] Grattafiori et al. The llama 3 herd of models. 2024.
- [Hart and Green, 2012] Herbert Lionel Adolphus Hart and Leslie Green. *The Concept of Law*. Oxford University Press, 2012.
- [Jolls et al., 1998] Christine Jolls, Cass Robert Sunstein, and Richard H. Thaler. A behavioral approach to law and economics. *Microeconomics: Decision-Making under Risk & Uncertainty eJournal*, 1998.
- [Kahneman and Tversky, 1979] Daniel Kahneman and Amos Tversky. Prospect theory: An analysis of decision under risk. 1979.
- [Kelly and Setman, 2020] Daniel Kelly and Stephen A. Setman. The psychology of normative cognition. 2020.

- [Kratochwil, 1989] Friedrich V. Kratochwil. *Rules, Norms, and Decisions: On the Conditions of Practical and Legal Reasoning in International Relations and Domestic Affairs*. Number 2 in Cambridge Studies in International Relations. Cambridge University Press, Cambridge, 1989.
- [Kutnick, 1986] Peter Kutnick. The relationship of moral judgment and moral action: Kohlberg’s theory, criticism and revision. In *Lawrence Kohlberg*, pages 133–156. Routledge, 1986.
- [Leng and Yuan, 2023] Yan Leng and Yuan Yuan. Do llm agents exhibit social behavior? *ArXiv*, abs/2312.15198, 2023.
- [Leslie et al., 2004] Alan M. Leslie, Ori Friedman, and Tim P. German. Core mechanisms in ‘theory of mind’. *Trends in Cognitive Sciences*, 8(12):528–533, 2004.
- [Li et al., 2024] Li et al. Culturepark: Boosting cross-cultural understanding in large language models. *Advances in Neural Information Processing Systems*, 37:65183–65216, 2024.
- [Madaan et al., 2023] Madaan et al. Self-refine: Iterative refinement with self-feedback. In *Thirty-seventh Conference on Neural Information Processing Systems*, 2023.
- [Mei et al., 2024] Mei et al. A turing test of whether ai chatbots are behaviorally similar to humans. *Proceedings of the National Academy of Sciences*, 121(9):e2313925121, 2024.
- [Meyer and Wieringa, 1994] John-Jules Ch. Meyer and Roel Wieringa. Deontic logic in computer science: normative system specification. 1994.
- [OpenAI, 2024] OpenAI. Gpt-4o system card. 2024.
- [Ostrom, 2000] E. Ostrom. Collective action and the evolution of social norms. *Journal of Natural Resources Policy Research*, 6:235 – 252, 2000.
- [Pan et al., 2022] Alexander Pan, Kush Bhatia, and Jacob Steinhardt. The effects of reward misspecification: Mapping and mitigating misaligned models. *ArXiv*, abs/2201.03544, 2022.
- [Park et al., 2022] Park et al. Social simulacra: Creating populated prototypes for social computing systems. *Proceedings of the 35th Annual ACM Symposium on User Interface Software and Technology*, 2022.
- [Park et al., 2023] Park et al. Generative agents: Interactive simulacra of human behavior. UIST ’23, New York, NY, USA, 2023. Association for Computing Machinery.
- [Piao et al., 2025] Piao et al. Agentsociety: Large-scale simulation of llm-driven generative agents advances understanding of human behaviors and society. *ArXiv*, abs/2502.08691, 2025.
- [Piquero et al., 2011] Piquero et al. Elaborating the individual difference component in deterrence theory. *Annual Review of Law and Social Science*, 7:335–360, 2011.
- [Posner, 1973] Richard A. Posner. *Economic Analysis of Law*. Little, Brown and Company, Boston, first edition, 1973.
- [Prelec, 1998] Drazen Prelec. The probability weighting function. *Econometrica*, 66:497–528, 1998.
- [Rao and Georgeff, 1995] Anand Srinivasa Rao and Michael P. Georgeff. Bdi agents: From theory to practice. In *International Conference on Multiagent Systems*, 1995.
- [Rawls, 2017] John Rawls. A theory of justice. In *Applied ethics*, pages 21–29. Routledge, 2017.
- [Raz, 1975] Joseph Raz. Practical reason and norms. 1975.
- [Richardson, 1994] Henry S Richardson. Practical reasoning about final ends. 1994.
- [Scherrer et al., 2023] Nino Scherrer, Claudia Shi, Amir Feder, and David M. Blei. Evaluating the moral beliefs encoded in llms. *ArXiv*, abs/2307.14324, 2023.
- [Sellin and Wolfgang, 1964] Thorsten Sellin and Marvin E Wolfgang. The measurement of delinquency. 1964.
- [Sharma et al., 2023] Sharma et al. Towards understanding sycophancy in language models. *ArXiv*, abs/2310.13548, 2023.
- [Shinn et al., 2023] Shinn et al. Reflexion: language agents with verbal reinforcement learning. In *Neural Information Processing Systems*, 2023.
- [Sripada and Stich, 2006] Chandra Sekhar Sripada and Stephen Stich. A framework for the psychology of norms. 2006.
- [Swisher and Dennison, 2016] Raymond R. Swisher and Christopher R. Dennison. Educational pathways and change in crime between adolescence and early adulthood. *Journal of Research in Crime and Delinquency*, 53:840 – 871, 2016.
- [Tversky and Kahneman, 1992] Amos Tversky and Daniel Kahneman. Advances in prospect theory: Cumulative representation of uncertainty. *Journal of Risk and Uncertainty*, 5:297–323, 1992.
- [Wei et al., 2022] Wei et al. Chain of thought prompting elicits reasoning in large language models. *ArXiv*, abs/2201.11903, 2022.
- [Wen et al., 2024] Wen et al. Language models learn to mislead humans via rlhf. *ArXiv*, abs/2409.12822, 2024.
- [Wu and Gonzalez, 1996] George Wu and Richard Gonzalez. Curvature of the probability weighting function. *Management Science*, 42:1676–1690, 1996.
- [Xie et al., 2024] Xie et al. Can large language model agents simulate human trust behaviors? *ArXiv*, abs/2402.04559, 2024.
- [Yang et al., 2024] Yang et al. Qwen2.5 technical report. *ArXiv*, abs/2412.15115, 2024.
- [Yax et al., 2023] Nicolas Yax, Hernan Anll’o, and Stefano Palminteri. Studying and improving reasoning in humans and machines. *Communications Psychology*, 2, 2023.
- [Young, 1996] H Peyton Young. The economics of convention. *Journal of economic perspectives*, 10(2):105–122, 1996.