

Can large language models overcome opportunistic behavior in the prisoner's dilemma?

Ting Kang

School of Economics and Management, Southwest Petroleum University, Chengdu 610500, China

How to cite this paper: Kang, T. (2025). Can large language models overcome opportunistic behavior in the prisoner's dilemma? *Advances in Engineering Research: Possibilities and Challenges*, 1(1), 152-165. ISSN Print: 3079-5192, ISSN Online: 3079-5206.

<https://doi.org/10.63313/AERpc.2005>

Published: 2025-03-31

Copyright © 2025 by author(s) and Erytis Publishing Limited.

This work is licensed under the Creative Commons Attribution International License (CC BY 4.0).

<http://creativecommons.org/licenses/by/4.0/>



Abstract

Large language models as a tool for simulating humans still need to be explored to a large extent. There is still a lack of exploration of whether the behavioral response of large language models to simple social stimuli is consistent with human behavior and whether it can behave in accordance with human behavior under certain environmental incentives. Testing their behavioral responses in classical behavioral game experiments provides a theoretical framework for evaluating the ability of LLMs to simulate humans. In this study, the behaviors of ERNIE-Speed-128K, qwen2-1.5b-instruct and gpt-4o-mini in the one-time prisoner's dilemma and the iterated prisoner's dilemma were studied, and the effects of three different motivational strategies on their behavioral responses in the iterated prisoner's dilemma were explored. The study found that the large language model can overcome the opportunistic behavior in the prisoner's dilemma through repeated interactions, and the large language model can perform better by adopting incentive strategies. This paper provides new insights for further research on the combination of large language models and social sciences.

Keywords

Large Language Models; The Prisoner's Dilemma; Opportunistic Behavior

1. Introduction

Large language models have received a lot of attention as an artificial intelligence designed to understand and generate human-like language. While they are gaining popularity as a workforce augmentation tool, the powerful emergent ability of large language models is expected to understand and mimic various human behaviors and preferences.

With the continuous development of large language models, there is a gradual increase in the trend of using large language models as agent-based human simulation tools in various fields, which provides new prospects for understanding the complex social interactions of human beings. Implicit in these studies is the important premise that large language models can simulate human responses. However, the ques-

tion of whether large language models can simulate human responses remains to be studied. In order to explore whether large language models can be used as human agents in the field of social sciences, it is necessary to study whether large language models can respond to simple social stimuli and whether they can be made to exhibit behaviors consistent with human responses through certain environmental incentives. Classical behavioral experiments provide an ideal basis for testing the responses of AI agents to archetypal social situations (Horton, 2023).

In this study, we focus on the behavioral responses of large language models in the prisoner's dilemma game. Firstly, the Folk theorem proves that game agents with selfish behaviors may produce cooperative behaviors in the process of repeated games. There have also been a series of experiments that have demonstrated that repetitive interactions can allow people to overcome opportunistic behaviors (Roth Alvin & Keith, 1978; Murnighan & Roth, 1983; Palfrey & Rosenthal, 1994; Engle-Warnick & Slonim, 2004, 2006a, 2006b; Feinberg & Husted, 1993; Fréchette & Sevgi, 2013; Dreber, Anna, Drew, Dreber, Drew, & Dreber, 2011). In this paper, we choose the one-time prisoner's dilemma game experiment and the infinitely repeated prisoner's dilemma game experiment to verify whether the large language model can overcome the opportunistic behavior in the prisoner's dilemma through repeated interaction. Through experiments, it is found that large language models can overcome opportunistic behaviors in prisoner's dilemma games through repetitive interactions. Then, aiming at the prisoner's dilemma caused by game subjects with selfish behavior, the influence of three different incentive strategies on the trust cooperative behavior of large language models in the process of repeated games is analyzed. Through experiments, it is found that three different incentive strategies can make the large language model better overcome opportunistic behavior, but different strategies acting on different large language models have different effects. In general, the strategy of tit for two tats is optimal, which is consistent with human experimental research.

The contribution of this article can be summarized as follows:

This paper examines LLMs at home and abroad, comprehensively studies whether LLM agents can simulate human behavior, and pays attention to the behavioral responses in the prisoner's dilemma.

It is found that LLM agents can overcome the opportunistic behavior in the prisoner's dilemma through repeated interactions, and at the same time, LLM agents can perform better by adopting incentive strategies. It lays a foundation for simulating complex human interactions, and provides a new idea for subsequent in-depth research on the analogy between large language models and humans.

2. Review of studies

The academic research on large language models as tools for simulating humans can be broadly divided into two categories. One is to study LLMs as human agents; The

other is to combine LLM with game theory for research.

LLMs as human agents: In the field of social sciences, more and more studies have studied LLMs as human agents, and the research on LLMs as human agents can be divided into simulations at the individual level and the social level. Used to simulate individual activities at the individual level, such as (Asfour & Murillo, 2023) explored the use of LLMs (specifically GPT-4) to simulate human responses to social engineering attacks, providing valuable insights to relevant researchers . (Dillion, Tandon, Gu, & Gray, 2023) uses LLMs as human agents for research in psychological science experiments . (Argyle, Busby, Fulda, Gubler, Rytting, & Wingate, 2023) experimented with giving LLMs a true socio-demographic background as human investigators. At the social level, LLMs are used to simulate social systems or social phenomena, such as (Park, O'Brien, Cai, Morris, Liang, & Bernstein, 2023), which has a town with 25 agents that users can interact with . All of these experiments are based on the premise that LLMs can simulate human behavior, but this premise has yet to be validated.

Combining LLMs with Game Theory: Research that combines LLMs with game theory can be divided into two categories. One is the ability to use game theory to explore LLMs, such as (Brookins, & DeBacker, 2024), who asked GPT-3.5 to play two classic games to understand LLMs' basic preferences for fairness and cooperation . (Mei, Xie, Yuan, & Jackson, 2024) developed a Turing test to assess the behavioral and personality traits exhibited by LLMs. The other is to study whether LLMs can replicate human behavior in existing game experiments, and the research in this field needs to be further explored, the most classic of which is that (Aher, Arriaga, & Kalai, 2023) designed a method to simulate Turing experiments to study whether LLMs can reproduce human behavior in classical experiments .

Early interdisciplinary research began to explore the application of LLMs in game theory, which has certain value, but there are still some limitations: 1) focusing on foreign LLMs, there is a lack of comparison of LLMs of different magnitudes; 2) Most of them focus on simple single-round games or simple iterative game experiments, and the abstraction ability of social prototype scenarios is limited; and 3) lack of research on the performance of LLMs under different incentives.

Therefore, this paper focuses on the prisoner's dilemma game in classical behavioral experiments, focusing on repetitive interactive behaviors. At the same time, the impact of different incentives on LLM behavior in repetitive interactions is studied, which provides a new idea for the follow-up in-depth study of whether LLM can simulate humans.

3. Background to the Prisoner's Dilemma

3.1. The one-time prisoner's dilemma

The one-time prisoner's dilemma is a classic game theory model used to study the behavior patterns of individuals in the face of choices between cooperation and be-

trayal. In this model, two actors must choose independently between the two actions and distribute the benefits based on the combination of the outcomes chosen by both. R: Reward for cooperation, T: Temptation for betrayal, S: Reward for stupidity, P: Punishment for betrayal. The classical structure of the game is defined by the payout levels $T > R > P > S$. Table 3-1 shows the income structure of the one-time prisoner's dilemma.

Table 3-1. One-time prisoner's dilemma benefit matrix.

	cooperate	betray
cooperate	R,R	S,T
betray	T,S	P,P

Players make decisions without future interactions, which leads to a seemingly paradoxical situation: although cooperation is beneficial to both parties, each individual tends to choose betrayal to maximize their own interests, resulting in a worse outcome for both parties.

3.2. The iterated prisoner's dilemma

The iterated prisoner's dilemma is a situation in which participants repeatedly play a strategic game G , which is called a component game. In this process, the set of behaviors of each participant is closely linked and the preference relationship of each participant is continuous. There is no limit to the number of times G can be performed, participants choose their actions at the same time on each occasion, and when actions are taken, one participant knows the actions that were previously chosen by all participants. In the iterated prisoner's dilemma, the component game is the one-time prisoner's dilemma, and on the basis of the one-time prisoner's dilemma payment level, the countermeasures cannot get out of the "dilemma" by betraying the other party in turn, that is, it is assumed that the "reward for the cooperation between the two parties" is greater than the average value of the "temptation to betray" and the "reward for the fool" (i.e., $R > (T+S)/2$).

For the classic "prisoner's dilemma" model, there is only one Nash equilibrium, where both sides choose the "betrayal" strategy at the same time. But if the game is infinitely repetitive, both sides will adopt a "cooperative" strategy, and the "cooperative" strategy becomes the only Nash equilibrium.

3.3. Infinitely repeat motivational strategies in the prisoner's dilemma

The infinite repetition of the prisoner's dilemma game not only provides a framework for understanding complex interactions in theory, but also provides valuable insights into solving the dilemma between cooperation and betrayal in practice, and has been studied in many literatures, the core idea of which is that repeated interactions can allow people to overcome opportunistic behavior.

Therefore, appropriate strategies can be used to motivate selfish large language models to participate in cooperation, and three typical incentive strategies in the infinitely repeated prisoner's dilemma are described in this section, and equilibrium

boundaries are obtained by inference.

The total profit of repeated games is the sum of the gains of the players at each stage, and the present value of the total profit is:

$$\pi = \sum_{t=1}^T \delta^{t-1} \pi_t (0 < \delta < 1, T \text{ can be } \infty) \quad (1)$$

In the Repeat Prisoner's Dilemma, the only reason for the existence of cooperation can only be that for any participant, the full benefit of choosing to betray at period t is not as good as the benefit of continuing to maintain cooperation at period t .

The GRIM strategy is to choose to cooperate in the first round, and then continue to cooperate as long as the opponent chooses not to betray. Once the other party betrays, then it will always choose to betray.

The total benefit for a player who chooses to betray in period t and in each subsequent episode is:

$$\begin{aligned} V_1 &= R \times (1 + \delta + \delta^2 + \dots + \delta^{t-2}) + T\delta^{t-1} + P \times (\delta^t + \delta^{t+1} + \dots) \\ &= R \times \frac{1 - \delta^{t-1}}{1 - \delta} + T\delta^{t-1} + P \times \frac{\delta^t}{1 - \delta} \end{aligned} \quad (2)$$

The total benefit for players who choose to cooperate in period t and choose to cooperate in each subsequent period is:

$$\begin{aligned} V_2 &= R \times (1 + \delta + \delta^2 + \dots + \delta^{t-2}) + R\delta^{t-1} + R \times (\delta^t + \delta^{t+1} + \dots) \\ &= R \times \frac{1 - \delta^{t-1}}{1 - \delta} + R\delta^{t-1} + R \times \frac{\delta^t}{1 - \delta} \end{aligned} \quad (3)$$

Cooperation can only be maintained when $V_2 > V_1$, that is, when $\delta > (T-R)/(T-P)$, cooperation can be reached.

The TFT strategy is to choose cooperation in the first round, and then use the strategy used by the opponent in the previous round in each subsequent round. In other words, if the opponent used cooperation in the previous turn, then the participants of the TFT strategy will choose to cooperate in this turn, and the opponent used betrayal in the previous round, and the participants of the TFT strategy will choose to betray in this turn.

The total benefit for a player who chooses to betray in period t and in each subsequent episode is:

$$\begin{aligned} V_1 &= R \times (1 + \delta + \delta^2 + \dots + \delta^{t-2}) + T\delta^{t-1} + P \times (\delta^t + \delta^{t+1} + \dots) \\ &= R \times \frac{1 - \delta^{t-1}}{1 - \delta} + T\delta^{t-1} + P \times \frac{\delta^t}{1 - \delta} \end{aligned} \quad (4)$$

The total benefit for players who choose to cooperate in period t and choose to cooperate in each subsequent period is:

$$\begin{aligned} V_2 &= R \times (1 + \delta + \delta^2 + \dots + \delta^{t-2}) + R\delta^{t-1} + R \times (\delta^t + \delta^{t+1} + \dots) \\ &= R \times \frac{1 - \delta^{t-1}}{1 - \delta} + R\delta^{t-1} + R \times \frac{\delta^t}{1 - \delta} \end{aligned} \quad (5)$$

Cooperation can only be maintained when $V_2 > V_1$, that is, when $\delta > (T-R)/(T-P)$, cooperation can be reached.

Tit for two tats is similar to TFT, but it only chooses to betray after the opponent betrays twice in a row. This makes players who use the TF2T strategy more tolerant. The total benefit for a player who chooses to betray in period t and in each subsequent episode is:

$$\begin{aligned} V_1 &= R \times (1 + \delta + \delta^2 + \dots + \delta^{t-2}) + T\delta^{t-1} + T\delta^t + P \times (\delta^{t+1} + \delta^{t+2} + \dots) \\ &= R \times \frac{1 - \delta^{t-1}}{1 - \delta} + T\delta^{t-1} + T\delta^t + P \times \frac{\delta^{t+1}}{1 - \delta} \end{aligned} \quad (6)$$

The total benefit for players who choose to cooperate in period t and choose to cooperate in each subsequent period is:

$$\begin{aligned} V_2 &= R \times (1 + \delta + \delta^2 + \dots + \delta^{t-2}) + R\delta^{t-1} + R \times (\delta^t + \delta^{t+1} + \dots) \\ &= R \times \frac{1 - \delta^{t-1}}{1 - \delta} + R\delta^{t-1} + R \times \frac{\delta^t}{1 - \delta} \end{aligned} \quad (7)$$

Cooperation can only be maintained when $V_2 > V_1$, that is, when $\delta > \sqrt{\frac{T-R}{T-P}}$, cooperation can be reached.

4. Research Methods

4.1. Subjects and experimental design

In this study, Baidu's ERNIE-Speed-128K, Alibaba's qwen2-1.5b-instruct, and OpenAI's gpt-4o-mini were selected as experimental subjects. This article accesses these models using an application programming interface (API) in python that allows users to interact with the models and collect their responses to specific prompts.

In this paper, two sets of experiments were conducted, the first set of experiments required three large language models to be conducted in the one-time prisoner's dilemma game and the repeated prisoner's dilemma game to investigate whether the large language model could overcome the opportunism in the one-time prisoner's dilemma game in repeated interactions. The second set of experiments requires three large language models to conduct experiments in the infinitely repeated prisoner's dilemma game under the benchmark experiment and three incentive strategies to investigate whether the large language model can better overcome the opportunism in the one-time prisoner's dilemma game under the three incentive strategies.

4.2. Experimental Procedure

Firstly, the corresponding experimental prompts were designed for the one-time prisoner's dilemma experiment, the repeated prisoner's dilemma game, and the infinitely repeated prisoner's dilemma game experiment under the three strategies, including the role of the experimental subject, the experimental rules, the experi-

mental object goal, and the experimental history. Secondly, different large language models are connected through the application programming interface, and the relevant experimental prompts are transmitted to different large language models, and the responses of the large language models are received. Finally, the received information is processed and analyzed, and the experimental results are obtained. The experimental procedure is shown in Figure 4-1.

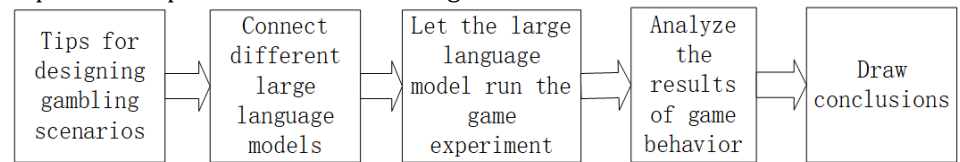


Figure 4-1. Schematic diagram of the experimental procedure.

5. Findings

5.1. Experimental manipulation

According to the theoretical basis and related inferences in Chapter 3, the return matrix of the single-round prisoner's dilemma is shown in Table 5-1 with $R=6$, $T=8$, $S=-4$, $P=-2$, and $\delta=0.9$.

Table 5-1. Single-round prisoner's dilemma game benefit matrix.

	cooperate	betray
cooperate	6,6	-4,8
betray	8,-4	-2,-2

In order to maximize the randomness of LLM performance, players who cooperate or betray will be randomly selected as opponents in the game with the same probability in each round.

- 1) One-time prisoner's dilemma game: Due to the randomness of the responses generated by LLM, this paper repeats $k=10$ times for each prisoner's dilemma game, and takes the average of all results as the final result of the experiment.
- 2) Repeated Prisoner's Dilemma Game: LLMs are asked to perform a series of repeated prisoner's dilemma games, each experiment containing $N=10$ rounds. Each experiment was repeated $k = 10$ times, and the average of all results from the first round was taken as the final result of the experiment.

By comparing the experimental results of the one-time repeated prisoner's dilemma game and the repeated prisoner's dilemma game, whether LLM can overcome opportunism in the prisoner's dilemma game through repeated interaction.

Have LLMs perform a series of repeated prisoner's dilemma games, with each experiment containing $N=10$ rounds. Each experiment was repeated $k = 10$ times, and the average of the experiment was taken as the final result of the experiment.

- 1) Benchmark experiment of infinitely repeated prisoner's dilemma game: In order to measure the comparison of LLM in the maximum random scenario and facing different incentive strategies as much as possible, the players who randomly selected to cooperate or betray as the opponent in the game with the same probability in

each round were selected as the benchmark experiment.

2) Infinite repetition of the prisoner's dilemma game experiment under different incentive strategies: The infinite repetition of the prisoner's dilemma game experiment under the cold strategy selects the player who adopts the cold strategy as the opponent for the experiment. The infinitely repeated prisoner's dilemma game experiment under the for-tat strategy will select the player who adopts the for-tat strategy as the opponent for the experiment. The infinitely repetitive prisoner's dilemma game experiment under the two-tooth strategy will select the player who takes the two-tooth strategy as the opponent for the experiment.

By comparing the results of the infinite repetition prisoner's dilemma game under different incentive strategies with the benchmark experiment of the infinite repetition prisoner's dilemma game, it is investigated whether LLMs can be motivated to cooperate by adopting appropriate incentive strategies to better overcome opportunism in the prisoner's dilemma game.

In order to realize these two sets of experiments, this paper designs corresponding prompts for each game experiment.

5.2. Experimental treatment

The three large language models are invoked through the API, and the corresponding game experiments are carried out with the matching opponents described in 5.1, and in the experimental process, the correct and complete response information is obtained by setting the recognition conditions for the response information, and the correct and complete response information is re-issued when the response error occurs, so as to facilitate the processing and research of the following text, and the response information of the obtained LLM is stored in the corresponding files respectively. Extract the results of the experiments in the file and convert the results of each round of each experiment into binary form, with 1 for cooperation and 0 for betrayal. The results of the experiment were analyzed by calculating the average cooperation rate.

In Experiment 1, the repeated prisoner's dilemma game experiment only counts the results of the first round. The i th result is A_i ($i \in k$), $A_i=1$ or 0. Its average cooperation rate is:

$$P = \frac{\sum_{i=1}^k A_i}{k}.$$

In Experiment 2, the result of the i th round j is A_{ij} ($i \in k, j \in N$), $A_{ij}=1$ or 0. The

average cooperation rate for round j is $C_j = \frac{\sum_{i=1}^k A_{ij}}{k}, j \in N$. The average cooperation

$$\text{rate in all rounds is } P = \frac{\sum_{j=1}^N C_j}{N}.$$

5. 3. Experimental treatment

By comparing the choice and average cooperation rate of different large language models in the first round of the one-time prisoner's dilemma game experiment and the repeated prisoner's dilemma game experiment, the corresponding scatter plots and bar charts were drawn.

As can be seen from Figures 5-1 and 5-2, the average cooperation rate of ERNIE-Speed-128K in the one-time prisoner's dilemma game is 20% in experiment 1, while the average cooperation rate in the repeated prisoner's dilemma game experiment increases to 100%, which makes ERNIE-Speed-128K more inclined to choose cooperation through repeated interactions.

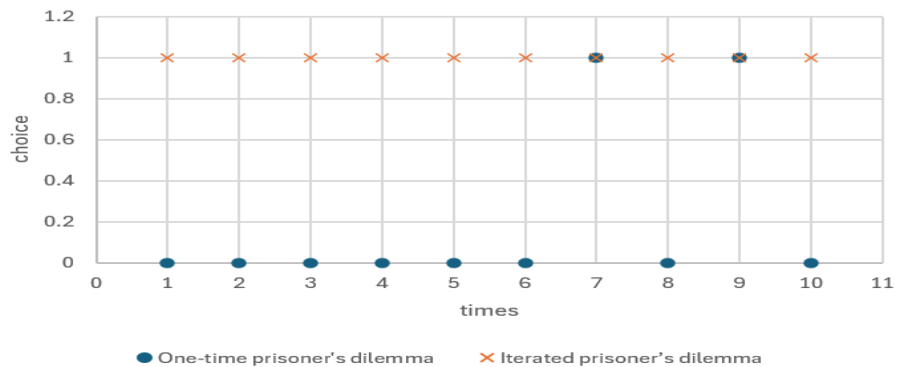


Figure 5-1. Experiment 1: Selection of ERNIE-Speed-128K.

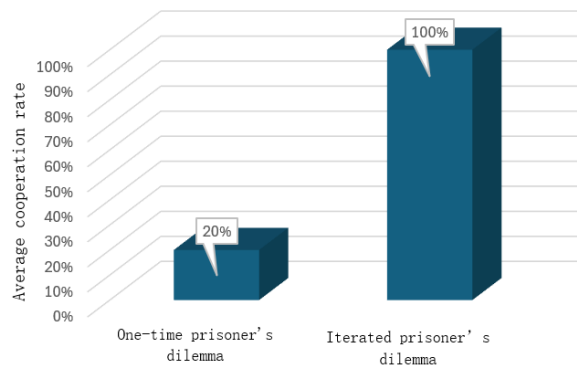


Figure 5-2. Experiment 1: Average cooperation rate of ERNIE-Speed-128K.

As can be seen from Figures 5-3 and 5-4, the average cooperation rate of qwen2-1.5b-instruct in the one-time prisoner's dilemma game is 20% in experiment 1, while the average cooperation rate in the repeated prisoner's dilemma game experiment increases by 90%.

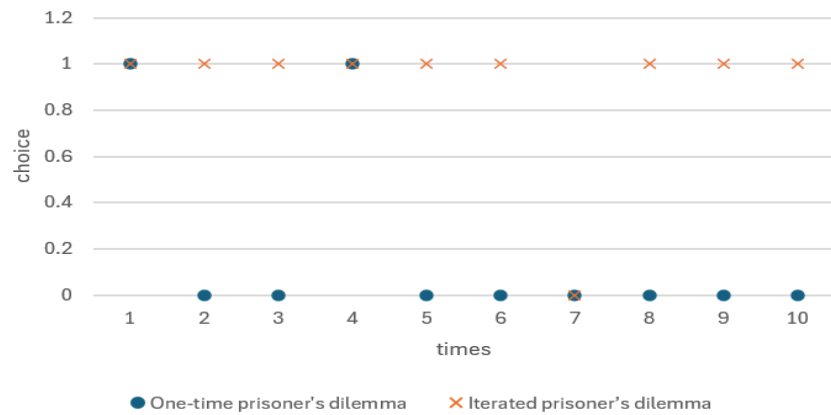


Figure 5-3. Experiment 1: Selection of QWEN2-1.5B-INSTRUCT.

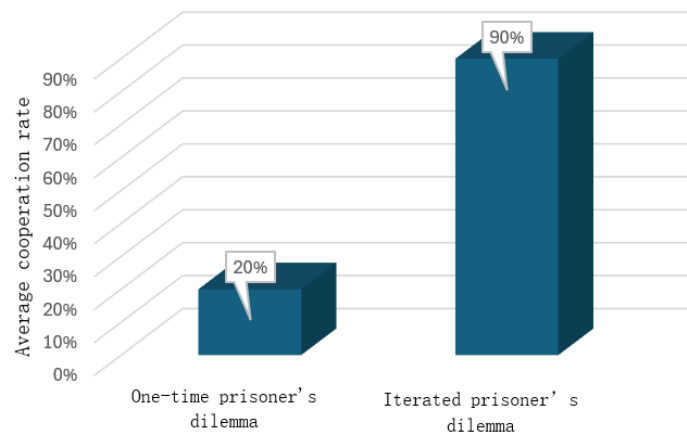


Figure 5-4. Experiment 1: The average cooperation rate of QWEN2-1.5B-INSTRUCT.

As can be seen from Figures 5-5 and 5-6, the average cooperation rate of GPT-40-mini in the one-time prisoner's dilemma game is 0% in experiment 1, while the average cooperation rate in the repeated prisoner's dilemma game experiment increases to 60%.

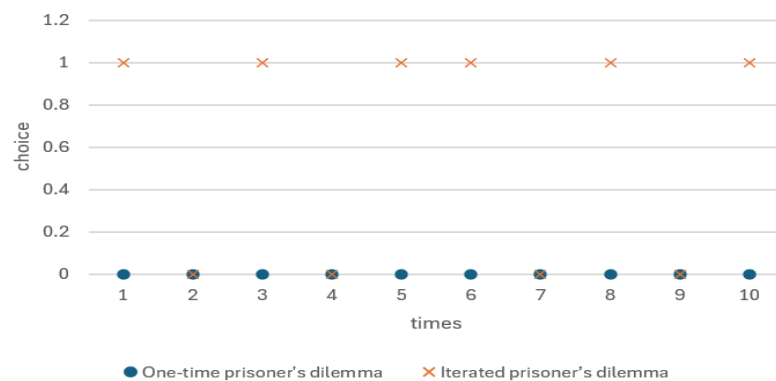


Figure 5-5. Experiment 1: Selection of GPT-40-mini.

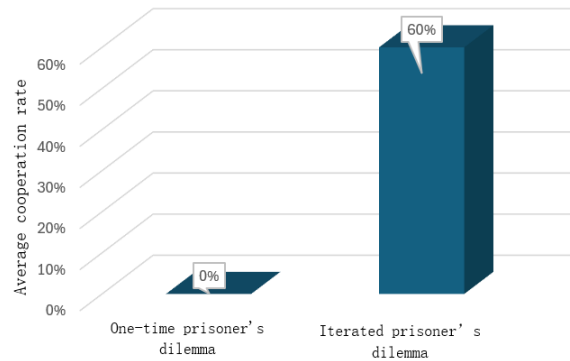


Figure 5-6. Experiment 1: Average cooperation rate of GPT-40-mini.

Through the experimental results of Experiment 1, it is found that the three models selected in this study can improve the average cooperation rate in the prisoner's dilemma game through repeated interaction, so the opportunism in the prisoner's dilemma can be overcome through repeated interaction.

According to the results of experiment 2 in 5.2, the performance of the large language model under different strategies can be plotted and compared with the benchmark experiment.

As can be seen from Figure 5-7, the probability of cooperation in the benchmark experiment is 92%, and the probability of cooperation in the repeated prisoner's dilemma game experiment with GRIM, TFT and TF2T is 100%. The experimental results show that the probability of large language model choosing cooperation can be improved by adopting different incentives, and the opportunistic behavior in the prisoner's dilemma can be better overcome than the random strategy. At the same time, the ERNIE-Speed-128K performs well in the face of all three incentives.

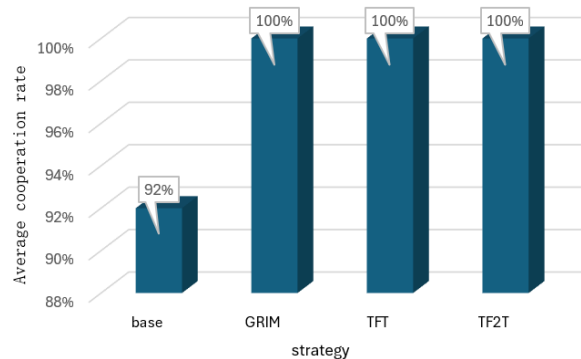


Figure 5-7. Average cooperation rate of ERNIE-Speed-128K under different strategies.

As can be seen from Figure 5-8, the probability of cooperation is 90% in the benchmark experiment, 97% in the infinite repetition prisoner's dilemma game experiment with the cold strategy, and 98% in the repeated prisoner's dilemma game experiment with TFT and TF2T. The experimental results show that the probability of

large language model choosing cooperation can be improved by adopting different incentives, and the opportunistic behavior in the prisoner's dilemma can be better overcome than the random strategy. And among the three incentives, TFT and TF2T have better results than the cold strategy.

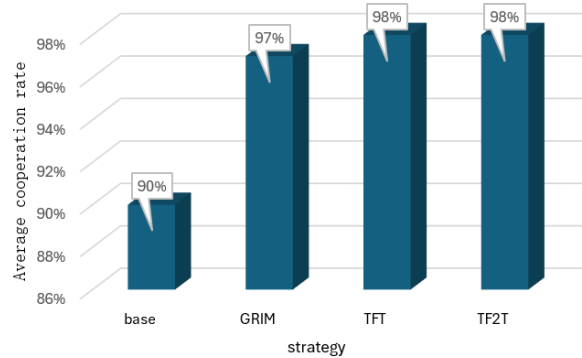


Figure 5-8. The average cooperation rate of qwen2-1.5b-instruct under different strategies.

As can be seen from Figure 5-9, the probability of cooperation is 26% in the benchmark experiment, 36% in the infinite repetition prisoner's dilemma game experiment with the cold strategy, 41% in the repeated prisoner's dilemma game experiment with TFT, and 47% in the repeated prisoner's dilemma game experiment with TF2T. The experimental results show that the probability of large language model choosing cooperation can be improved by adopting different incentives, and the opportunistic behavior in the prisoner's dilemma can be better overcome than the random strategy. And among the three incentives, the TF2T outperformed TFT and GRIM.

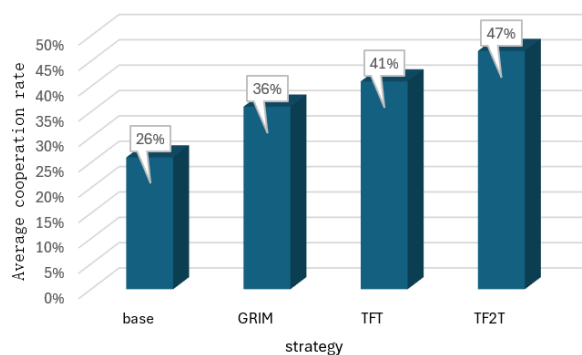


Figure 5-9. Average cooperation rate of GPT-4o-mini under different strategies.

Overall, in repeated interactions, the three large language models, ERNIE-Speed-128K, qwen2-1.5b-instruct, and gpt-4o-mini, all performed better when faced with different incentive strategies compared to the condition without any incentives. However, the effects of the three incentive strategies were not the same. For ERNIE-Speed-128K, all three strategies led it to choose full cooperation.

For qwen2-1.5b-instruct, the TFT and TF2T had similar effects. For gpt-4o-mini, the effects of the three strategies were all somewhat different.

6. Discussion and conclusions

This paper enriches the literature on the simulation of human behavior by large language models and provides a new research direction: whether certain incentive measures can improve the performance of large language models. Specifically, this paper examines the behavioral responses of three large language models, ERNIE-Speed-128K, qwen2-1.5b-instruct, and gpt-4o-mini, in the classic social interaction game of the Prisoner's Dilemma, using frequency to measure the results, which can to a certain extent truly reflect the most common behaviors of LLMs and reduce errors. When facing a one-time Prisoner's Dilemma, as selfish actors, LLMs will mostly choose to defect, which is consistent with the assumption of "rational people" in economics. At the same time, through repeated interactions, the cooperation rate of LLMs will increase, which is consistent with the results of early behavioral experiments involving LLMs [15-16], further proving that LLMs are consistent with human behavior, that is, they can overcome opportunistic behavior in the Prisoner's Dilemma through repeated interactions. Based on this experiment, this paper focuses on the impact of three strategies, namely tit-for-tat, tit for two tats, and grim trigger, on the behavior of LLMs in repeated interactions. The study shows that these three incentive strategies can all increase the cooperation rate of LLMs and better overcome opportunistic behavior in the Prisoner's Dilemma compared to the baseline experiment. Among these three strategies, the tit for two tats incentive strategy enables LLMs to perform the best, which is consistent with the results of human experiments.

Through this research, it further demonstrates that LLMs can simulate human behavior in classic social interaction experiments and preliminarily indicates that certain incentive measures can improve the performance of LLMs. However, the current research still has certain limitations: 1) The application of management theory to the simulation of human behavior by LLMs in this paper is only a preliminary exploration, and more in-depth and extensive experiments are needed for further verification; 2) This paper only adopts a fixed payoff structure and discount factor. Subsequent extensive research on the behavioral performance under different payoff structure scenarios is necessary, which will help to discover whether LLMs can also simulate human behavior well in more flexible environments. In future research, more in-depth analysis and discussion are still needed.

References

- [1] Horton, J. J. (2023). Large language models as simulated economic agents: What can we learn from Homo silicus?
- [2] Roth, Alvin, J. Keith Murnighan. (1978). Equilibrium behavior and repeated play in prisoners' dilemma games. *Journal of Mathematical Psychology*, 2, 189-198.

- [3] Murnighan, J.K., & Roth, A.E. (1983). Expecting Continued Play in Prisoner's Dilemma Games. *Journal of Conflict Resolution*, 27, 279 - 300.
- [4] Palfrey, T.R., & Rosenthal, H. (1994). Repeated Play, Cooperation and Coordination: An Experimental Study. *The Review of Economic Studies*, 61, 545-565.
- [5] Engle-Warnick, J., & Slonim, R. L. (2004). The evolution of strategies in a repeated trust game. *Journal of Economic Behavior & Organization*, 55(4), 553-573.
- [6] Engle-Warnick, J., & Slonim, R.L. (2006a). Learning to trust in indefinitely repeated games. *Games Econ. Behav.*, 54, 95-114.
- [7] Engle-Warnick, J., & Slonim, R.L. (2006b). Inferring repeated-game strategies from actions: evidence from trust game experiments. *Economic Theory*, 28, 603-632.
- [8] Feinberg, R.M., & Husted, T.A. (1993). An Experimental Test of Discount-Rate Effects on Collusive Behaviour in Duopoly Markets. *Journal of Industrial Economics*, 41, 153-160.
- [9] Fréchette, G. R., & Yuksel, S. (2013). Infinitely repeated games in the laboratory: Four perspectives on discounting and random termination. Working Paper.
- [10] Dreber, A., Fudenberg, D., & Rand, D. G. (2011). Who cooperates in repeated games? Working Paper.
- [11] Asfour, M. & Murillo, J. C. (2023). Harnessing large language models to simulate realistic human responses to social engineering attacks: A case study . *International Journal of Cyber-security ntelligence & Cybercrime*, 6(2), 21-49.
- [12] Dillion, D., Tandon, N., Gu, Y., & Gray, K. (2023). Can AI language models replace human participants?. *Trends in cognitive sciences*, 27(7), 597-600.
- [13] Argyle, L.P., Busby, E., Fulda, N., Gubler, J.R., Rytting, C., & Wingate, D. (2022). Out of One, Many: Using Language Models to Simulate Human Samples. *Political Analysis*, 31, 337 - 351.
- [14] Park, J. S., O'Brien, J., Cai, C. J., Ringel Morris, M., Liang, P., & Bernstein, M. S. (2023). Generative agents: Interactive simulacra of human behavior. In *Proceedings of the 36th Annual ACM Symposium on User Interface Software and Technology (UIST '23)* (Article 2, pp. 1-22). Association for Computing Machinery.
- [15] Brookins, P., & DeBacker, J. (2024). Playing games with GPT: What can we learn about a large language model from canonical strategic games? *Economics Bulletin*, 44(1), 25-37.
- [16] Mei, Q., Xie, Y., Yuan, W., & Jackson, M. O. (2024). A Turing test of whether AI chatbots are behaviorally similar to humans. *Proceedings of the National Academy of Sciences of the United States of America*, 121(9), e2313925121. <https://doi.org/10.1073/pnas.2313925121>
- [17] Aher, G., Arriaga, R. I., & Kalai, A. T. (2023). Using large language models to simulate multiple humans and replicate human subject studies. In *Proceedings of the 40th International Conference on Machine Learning (ICML'23)* (Vol. 202, Article 17, pp. 337-371). JMLR.org.