

# Exploring the Choice Behavior of Large Language Models

Anonymous ACL submission

## Abstract

Large Language Models (LLMs) are increasingly deployed as human assistants across various domains where they help to make choices. However, the mechanisms behind LLMs' choice behavior remain unclear, posing risks in safety-critical situations. Inspired by the intrinsic and extrinsic motivation framework within the classic human behavioral model of Self-Determination Theory and its established research methodologies, we investigate the factors influencing LLMs' choice behavior by constructing a virtual QA platform that includes three different experimental conditions, with four models from GPT and Llama series participating in repeated experiments. Our findings indicate that LLMs' behavior is influenced not only by intrinsic attention bias but also by extrinsic social influence, exhibiting patterns similar to the Matthew effect and Conformity. We distinguish independent pathways of these two factors in LLMs' behavior by self-report. This work provides new insights into understanding LLMs' behavioral patterns, exploring their human-like characteristics.

## 1 Introduction

Large Language Models (LLMs) are increasingly being adopted across numerous domains and often encounter practical scenarios where a choice needs to be made. For example, recommending some books for users without any explicit user preferences (He et al., 2023), analyzing open questions that different cultures have different viewpoints (Li et al., 2024; Tao et al., 2024). Despite the growing reliance on LLMs in these scenarios, the mechanisms behind their choice behavior remain unclear, raising questions about how LLMs make their choices and the influencing factors behind. Therefore, uncovering their behavioral patterns and discovering the influencing factors not only advances research in LLMs' explainability and human likeness to help better understand the

behavioral patterns of LLMs but also offers new insights into identifying behavioral risks such as neglect and harmful behavior caused by bias or bad output of violating ethical standards.

Psychology has long served as a powerful tool in uncovering the intricacies of human cognition and behavior (Festinger and Katz, 1953; Edwards, 1954), such as scenario simulation (Zimbardo et al., 1971), neuroscience (Uttal, 2011) and psychometrics (Furr, 2021). In recent years, the growing complexity and interpretability challenges of LLMs have spurred interdisciplinary approaches from artificial intelligence and psychology. By leveraging psychological methodologies, researchers are gaining deeper insights into the human-like behavioral characteristics and underlying mechanisms exhibited by these models (Shiffrin and Mitchell, 2023; Burnell et al., 2023; Hagendorff et al., 2024). In terms of evaluation, psychometric insight has made it possible to assess human-like psychological traits in LLMs (Wang et al., 2023), such as personality (Serapio-García et al., 2023), theory of mind (Strachan et al., 2024). In terms of evaluation eliciting capabilities, advances in psychological research on reasoning, emotion, and motivation have enabled improvements in response quality through techniques such as the generation of multiple chains of thought (Zhang et al., 2022).

Serving as a fundamental theory of human behavior, Self-Determination Theory (SDT) distinguishes intrinsic and extrinsic motivation (Deci and Ryan, 2013, 2000). The intrinsic motivation focuses on exploratory, playful, and curiosity-driven behaviors while extrinsic motivation focuses on instrumental value (Ryan and Deci, 2000).

In this paper, we analyze the choice behavior of LLMs from the perspective of SDT. We first propose three questions: (1) Will LLMs also focus too much on one part of the options and ignore the other due to the influence of intrinsic factors when faced with choices? We use attention bias to

describe this situation. (2) Will LLMs change their choice behavior to some extent due to extrinsic factors when faced with choices? Social influence is one of the important sources of extrinsic influence. (3) If both intrinsic and extrinsic factors are present, how do they interact with each other?

To address these questions, we design an experiment platform inspired by (Salganik et al., 2006), which explored social influence patterns through virtual music web pages. We begin by collecting two question sets spanning various domains, such as education, culture and ecology. These questions are then presented on a virtual QA platform modeled after Quora (Quo), where LLMs can view, like, and answer them.

On this QA platform, We employ a controlled variable approach to observe the choice behaviors of four models from GPT and Llama series across two distinct question sets under three experimental conditions. Each experiment is replicated three times to ensure the reliability of the results.

After analyzing the results, our key findings are:

- **Intrinsic Bias in LLMs’ Choice-Making :** LLMs exhibit internal biases, choosing certain topics like science or technology over others, similar to how humans have personal preferences.
- **Consistent Social Influence Reinforces Accumulation of Prior Behavior :** LLMs, like humans, are influenced by popularity. Topics with more views tend to receive more attention, reinforcing a bias toward the already popular, resembling the Matthew Effect(Rigney, 2010).
- **Conflicting Social Influence Leads to Behavioral Shifts :** When social influence is manipulated and conflicts with bias, LLMs shift their attention to a certain extent from intrinsic data-driven to socially influenced directions, resembling the Conformity Effect(Bernheim, 1994).
- **Distinct Pathways of Intrinsic and Extrinsic Dimensions :** Intrinsic and extrinsic influences affect LLMs independently, with separate mechanisms guiding how they balance personal biases with external social signals.

To summarize, our contributions are three-fold: (1) We pioneer applying self-determination theory to analyze decision-making mechanisms of LLMs. (2) We constructed a virtual QA platform for observation, which simultaneously ensures authenticity

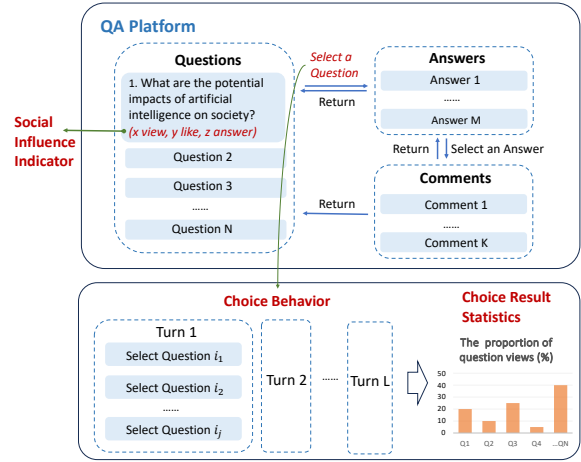


Figure 1: QA Platform

and controllability. (3) We demonstrate how intrinsic biases and extrinsic social influence jointly drive LLMs’ choice behavior.

## 2 Method

### 2.1 QA platform

To simulate real-world scenarios while allowing controlled manipulation of experimental conditions (e.g., question placement, presence of social information), we construct a QA platform modeled after Quora. As shown in Figure 1, we present our questions on platform where LLMs can view, like, answer and comment. Each question is presented to LLMs accompanied by views, answers, and likes. LLMs can like some questions and select a question to explore in depth. Once a question is selected, all answers related to the chosen question, along with their views, comments, and likes, are presented to LLMs. LLMs can choose to answer the question, like existing answers, and then select one answer to view its comments. LLMs can also choose to like the comments. Each of the aforementioned actions contributes to the corresponding item’s count. For instance, viewing a question increases the question’s view count by 1. Detailed interaction process and the prompt settings for the process are in Appendix A.1 and A.2.

Each experiment consists of multiple independent turns. In a turn, the LLM interacts with the QA platform, viewing multiple questions until it chooses to end the turn. Once a turn ends, the context is cleared, ensuring there is no shared context between different turns, thus maintaining complete independence of choice-making across turns.

To address the potential issue of zero-value data

points skewing the experimental results, we use *views* as the primary analysis metric, as it always has higher values than *likes* and *answers*

## 2.2 Choice Behavior and Two Factors

We define the concepts of choice behavior, attention bias of LLMs, and social influence within our QA platform.

**Choice Behavior** The choices LLMs make when faced with multiple question candidates. The final distribution of question views after multi-turn interactions reflects the cumulative outcomes of these choices, which will be the core of our analysis.

**Social Influence** The existing choice results when LLMs make choices. On our platform, the metrics (view, like, answer), which are appended to the content of the questions, serve as indicators of social influence. For example, if the indicator values for Question 1 are (10, 6, 4), it means that before the LLMs make a choice, Question 1 has been viewed 10 times, liked 6 times, and received 4 answers.

**Attention Bias** A classic paper of self determination (Deci, 1971) proposes that when there are no extrinsic reasons to perform a task (e.g., no rewards or approval), the longer an individual engages in a task, the stronger their intrinsic motivation for it. Measuring time spent by LLMs is challenging, we define attention bias in our experiment as the frequency of repeated selection among certain options without social influences.

## 2.3 Three Conditions

First, we divide the experimental conditions into two categories based on whether the social influence indicator is visible to the LLMs: (1) independent condition and (2) social influence condition (SI). Based on the SI condition, we then introduce a new condition: (3) induced social influence condition (Induced-SI).

**Independent Condition** The social influence indicators, including the number of views, answers, and likes that each question receives, are not visible to LLMs. The condition allows us to measure their intrinsic bias without external influence.

**Social influence Condition** The social influence indicator is visible to LLMs when they make choices. When LLMs interact in the platform, interaction behavior is also updated to each SI indicator

synchronously. By default, the initial values of all questions are set to 0. This design allows us to examine how LLMs make choices as social influence begins to accumulate from a neutral starting point.

**Induced social influence Condition** Building on the SI condition, we introduce the Induced SI condition to explore scenarios where certain questions are intentionally given an advantage in terms of social influence. Specifically, we set the initial values of specific questions to non-zero values. This allows us to investigate how varying levels of social influence impact the LLMs' choice-making and attention allocation.

## 2.4 Experiment Setup

**Question Sets** We obtain 2 question sets and put them on the QA platform as the basis for interaction. Table 1 shows the question sets used in our research. More details are shown in Appendix C.

**Selected Models** We select a total of 4 models, including two proprietary models (GPT-4-1106-preview (OpenAI et al., 2024), GPT-4o-2024-05-13 (OpenAI)) and two open-source models (Llama 3.1: 70B (Grattafiori et al., 2024) and Llama 3.3: 70B (Meta AI)). This choice aims to maintain diversity and representativeness within the constraints of limited resources.

**Repeated Experiments** To mitigate the impact of randomness and ensure reliability, each experimental setting is repeated three times.

**Rating and Shuffle** To eliminate the effect of question order in the context, we shuffle the orders of questions and show them as a random sequence each time the LLMs make choices. Additionally, to prevent simple reaction patterns, such as choosing questions in sequential order (e.g. 1, 2, 3...), we require the LLMs to assign a score to each question according to "comprehensive aspects".

## 3 Experiments

We conduct experiments under three conditions. In the independent condition, we focus on the attention bias (3.1). Under the SI condition, we explore how social influence affects LLMs' choice behavior (3.2.1). Under the Induced-SI condition, we examine the choice behavior of LLMs when the intrinsic bias and extrinsic social influence are in conflict

<sup>1</sup> See details in <https://www.science.org/content/resource/125-questions-exploration-and-discovery>

Table 1: Question sets and acquisition process

| Question set | Explanation                                                                                                                    |
|--------------|--------------------------------------------------------------------------------------------------------------------------------|
| G            | GPT4 generated questions, including question types and question content, and is required to be as comprehensive as possible.   |
| S            | 36 questions randomly selected from <i>125 QUESTIONS: EXPLORATION AND DISCOVERY</i> published by <i>Science</i> . <sup>1</sup> |

(3.2.2). Finally, we investigate the significance of attention bias and social influence on behavioral pattern of LLMs at the mechanistic level (3.3).

### 3.1 Attention Bias of LLMs

To assess attention bias in the four models across question sets G and S, we conduct experiments under independent conditions. Each model completed three experimental repetitions, each consisting of 100 turns. We calculate the proportion of times each question is selected relative to the total selection times, mapping the LLMs’ choice distribution across the entire question set.

#### 3.1.1 Bias in G Set

To begin, question set G comprehensively covers various topics, including 12 topics, with 5 questions per topic, resulting in a total of 60 questions.

We rearrange the choice distribution in descending order of the percentage of views and draw a Pareto chart. The experimental results in the G question set are shown in Figure 2. The results show that there are significant differences in attention allocation of LLMs among different questions, and the comparison between popular and unpopular questions is clearly reflected.

Our conclusion implies meaningful research prospects, that is, LLMs have attention bias among different options when facing choices, and they will always pay attention to some of them, but if we ask LLMs these questions one by one, they will try their best to answer each question. This shows that our method is an effective method that can reveal the behavior pattern of LLMs.

Question-level attention bias is hard to interpret directly. Thus, we analyze topic-level bias by aggregating the views of all questions within each topic to determine the total view count per topic. Then we find a clear relationship between LLMs’

attention bias and topics. Results for Set G are shown in Figure 3. For comparison, we highlight the top three and bottom three topics by view proportion from the four models’ experiments. All topics’ distribution is available in Appendix D.1.

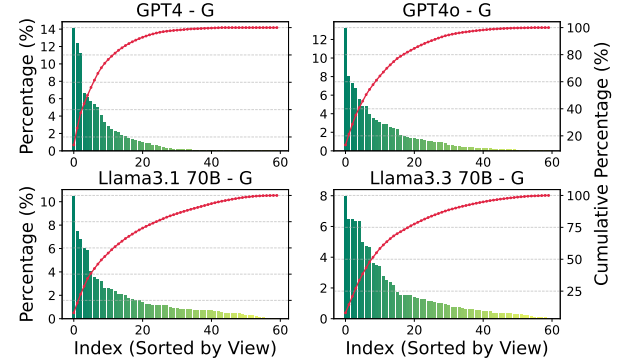


Figure 2: Views distribution in G set (The curve represents the cumulative process of the number of views, that is, the percentage of the sum of views in the total increases with the X axis)

It is evident that all models exhibit unequal distribution of attention in various topics. Interestingly, all models also have certain commonalities in the most popular and least unpopular topics. *Technology* and *Science* are most popular topics across all models. *History* and *Society*, *Arts* and *Culture* are all unpopular in the selection of at least three models. Meanwhile, the GPT series models do not focus on *Sport and Recreation*, while the Llama series models do not focus on *Food and Agriculture*.

#### 3.1.2 Bias in S Set

To expand our conclusions, we conduct further experiments on the most popular topics across four LLMs. Set S consists of 36 questions, including 12 topics, with 3 questions per topic. All topics are related to science and technology. The results are shown in Figure 4 and Figure 5.

First, we discover that, similar to the situation in the G set at the question level, there are significant differences in attention among questions for LLMs in set S. This indicates that attention bias is also present in set S. Next, from a topic-level perspective, we find that there are common patterns in set S as well. For instance, all four models show a strong attention bias for *Neuroscience*, whereas *Math* receives more attention from three models. At the same time, *Energy Science* and *Biology* tend to be less favored.

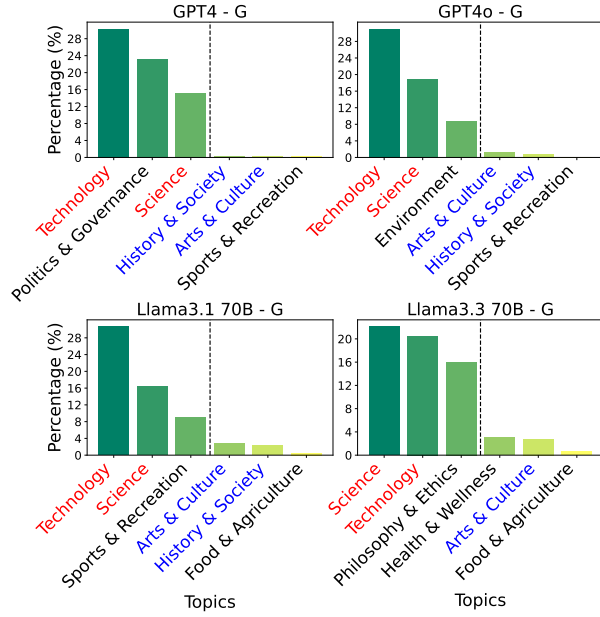


Figure 3: View proportion of top three and bottom three topics in G set (Topics in red/blue consistently ranked top/bottom three in  $\geq 3$  out of 4 models)

### 3.2 Behavior Patterns Under Social Influence

To examine how social influence impacts LLMs' choice behaviors, we implemented two experimental protocols: In SI condition, we initialize all questions with zero-valued SI indicators to equalize social influence level (3.2.1); In induced SI condition, unpopular questions (lowest selection rate in independent condition) are artificially assigned non-zero SI values, thereby creating controlled conflict between attention bias and social influence (3.2.2).

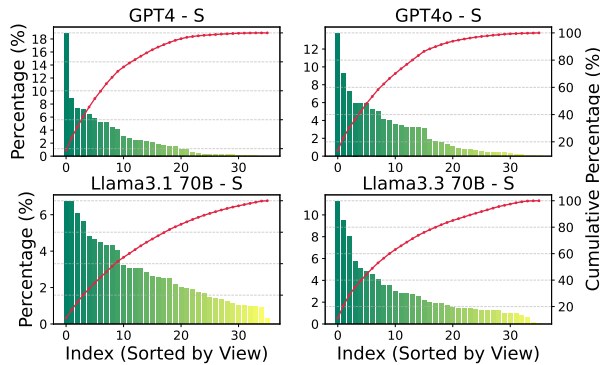


Figure 4: Views distribution in S set (The curve represents the cumulative process of the number of views, that is, the percentage of the sum of views in the total increases with the X axis)

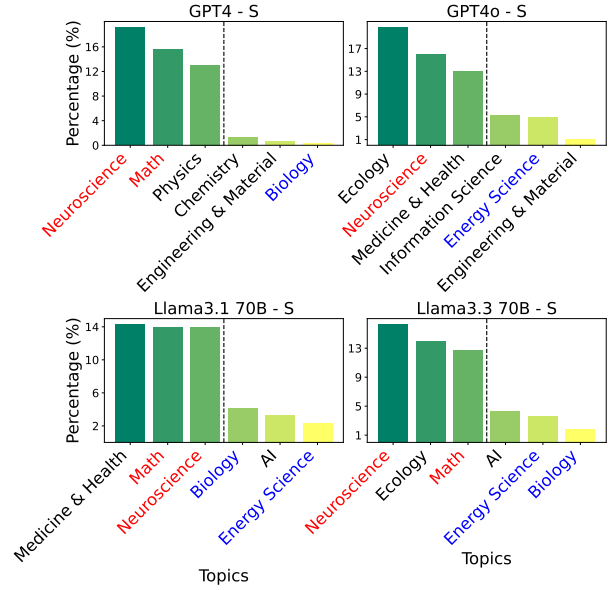


Figure 5: View proportion of top three and bottom three topics in S set (Topics in red/blue consistently ranked top/bottom three in  $\geq 3$  out of 4 models)

#### 3.2.1 Social influence makes LLMs reinforce their biased choices

**Social influence makes LLMs' choices more concentrated** The first observation is that all LLMs' choices become more concentrated in SI condition compared to independent condition. Figure 6 compares the LLMs' choice inequality measured by the Gini coefficient of question views between two conditions. The results show that all LLMs' choices become more unequal, which means more choices are concentrated on fewer questions, aligning well with the Conformity Effect where individuals follow majority-preferred options under social influence.

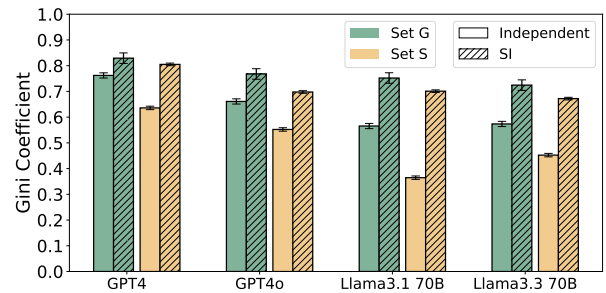


Figure 6: Gini coefficient of the distribution of views under independent condition and SI condition.

**LLMs' choices still align well with their attention bias** When social influence makes LLMs' choices more concentrated, we explore whether the

LLMs’ choices still align with their attention bias. We evaluate it by measuring the Spearman rank correlation coefficients of question views between independent condition and SI condition. The results in Table 2 show that the coefficients are high in all situations, indicating a strong similarity in the ranking order of questions between two conditions. This suggests that in the SI condition, the LLMs’ choices still align well with their attention bias. We also present the results in Figure 7 to show the conclusion visually.

| Models       | Set G           | Set S           |
|--------------|-----------------|-----------------|
| GPT4         | $0.69 \pm 0.03$ | $0.71 \pm 0.05$ |
| GPT4o        | $0.80 \pm 0.02$ | $0.84 \pm 0.02$ |
| Llama3.1 70B | $0.59 \pm 0.08$ | $0.67 \pm 0.06$ |
| Llama3.3 70B | $0.72 \pm 0.05$ | $0.66 \pm 0.04$ |

Table 2: Spearman’s rank correlation coefficient of views distribution under independent condition and SI condition.

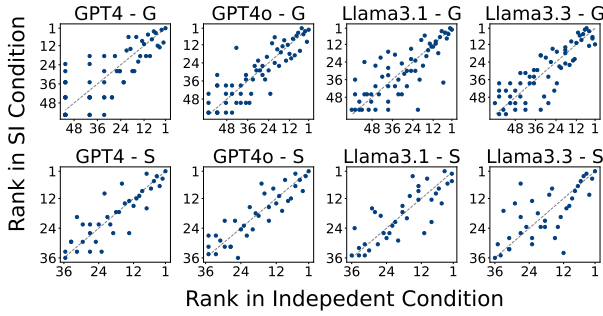


Figure 7: Rank of question views between independent condition and SI condition (Each point is a question. X-axis and Y-axis are rankings of views under independent condition and SI condition respectively)

**Matthew’s Effect** Based on the above results, we find that social influence makes LLMs’ choices more concentrated on their attention bias, consistent with the Matthew effect. The mechanism works as follows: Initially, all question views are 0. LLMs make choices based on their attention bias. The views on these chosen questions produce a positive conformity effect to help them build initial advantage. Subsequently, LLMs are more likely to choose these questions and this will further expand their advantage. This is a positive feedback that makes LLMs reinforce their biased choice.

### 3.2.2 Behavior Shift under Conflict Situation

In this section, we explore the LLMs’ choice behavior when social influence conflicts with their

attention bias. We first select 9 **most unpopular** questions. Then, we manually set the initial view of each of them to 50 (half of the total number of turns in independent condition experiments), while keeping the initial view of all other questions at 0. For simplicity, we refer to these questions as "induced questions". Finally, we investigate whether these questions would be selected in the context of high-level social influence.

As shown in Figure 8, we compare the view percentage of induced questions under independent condition and induced SI condition. Following artificial intervention, the selection rate of these induced questions increases significantly, with some even becoming very popular. Notably, most popular questions in independent condition still maintained their high selection rates under conflicting conditions. This pattern demonstrates that social influence can partially override LLMs’ inherent attention biases when the two factors conflict, while models still take both factors into account when making choices.

## 3.3 Research on Independent Pathways

Another method of assessing intrinsic motivation involves the use of self-reports that capture interest and enjoyment derived from the activity itself (Ryan, 1982; Harackiewicz, 1979). We draw inspiration from this methodology and use self-reporting method to explore the tendency of LLMs to consider intrinsic and extrinsic factors when making choices.

### 3.3.1 Experimental design

We conduct a self-report approach to further demonstrate our conclusions. Attention bias and social influence are well aligned with the definitions of intrinsic motivation and extrinsic motivation in self-determination theory. Therefore, six indicators about bias or social influence that LLMs will take into consideration when making question selection are generated to represent the impact sources of the intrinsic and extrinsic dimensions. Each indicator is a scoring standard, which is used to quantify the size of internal and external factors that LLMs consider when viewing questions. GPT-4, GPT-4o and Llama3.3: 70b<sup>1</sup> score each question of G and S on these six indicators on a scale of 1 to 100. Among the six indicators, Personal Goals or

<sup>1</sup>Llama3.1:70b was not included in the experiment because it refused to score on human dimensions such as interest or output the same score for all questions on indicators.

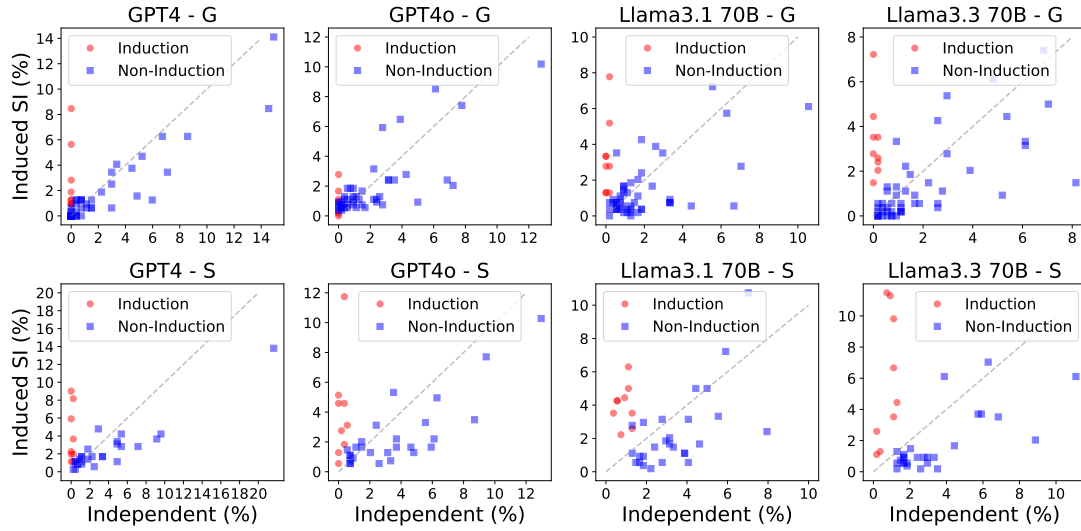


Figure 8: Percentage of views (subtract the initial number of views) under independent condition and induced social influence condition (Each point is a question. Red indicates artificially induced questions. X-axis and Y-axis are the percentage of views in the total views under independent condition and induced SI condition respectively)

Interests, Insight or Epiphany and Cognitive Dissonance represent the sources of bias’s influence, namely, intrinsic factors, while Expert Opportunity, Community Engagement and Interdisciplinary Connections represent the sources of social influence, namely, extrinsic factors. The explanations of indicators and prompts used scoring are detailed in the Appendix B.2.

### 3.3.2 Result

We calculate the Spearman correlation coefficients between the indicator ratings for each question and the number of views for each question in previous simulation experiments, along with the significance levels. The results are shown in Table 3.

Clearly, our findings capture the significant correlations between intrinsic and extrinsic factors in question choice behavior. We further conduct a factor analysis on the views and all indicators to explore whether the indicators are also mapped onto independent latent factors. Factor analysis is a statistical method used to explore the potential structure behind the observed variables and identify potential factors that cannot be observed directly. It can help us reduce the dimension of variables, reduce multiple related variables to a few core factors, and thus reveal the essential characteristics and internal relations of data.

The results are shown in Figure 9. Except for the results of Llama3.3 on the S question set, which already exhibit a clear factor structure, we perform orthogonal rotation, a method which makes each

variable more clearly belong to a specific factor while maintaining the independence between factors, on the factor loadings for the other results.

Our findings clearly identify two latent factor structures, with indicators loading onto separate factors. These results demonstrate that both intrinsic and extrinsic factors significantly influence LLMs’ behavior, and their influence pathways are relatively independent, aligning with the dual motivational driving model proposed in Self-Determination Theory.

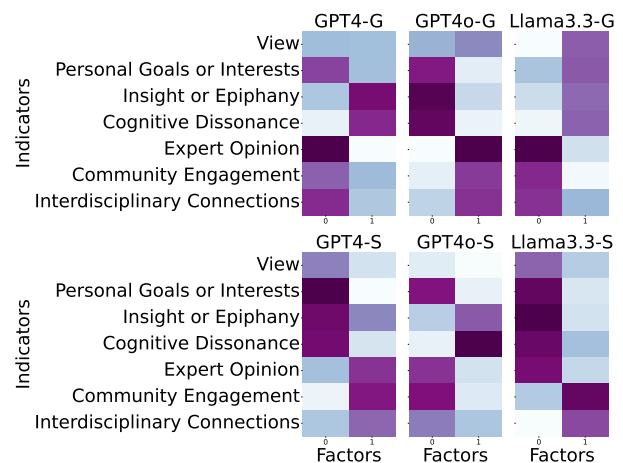


Figure 9: Indicators’ load on two potential dimensions(The X-axis is two potential factors obtained by factor analysis. Color depth indicates the load of the indicator on the potential factors.)

Table 3: Correlation between views and scoring of indicators(\*:  $p < 0.05$ , \*\*:  $p < 0.01$ )

| Evaluation dimensions         | Correlation to view |               |                  |              |               |                  |
|-------------------------------|---------------------|---------------|------------------|--------------|---------------|------------------|
|                               | GPT4<br>IN G        | GPT4o IN<br>G | Llama3.3<br>IN G | GPT4 IN<br>S | GPT4o IN<br>S | Llama3.3<br>IN S |
| Personal Goals or Interests   | 0.56**              | 0.57**        | 0.46**           | 0.80**       | 0.64**        | 0.24             |
| Insight or Epiphany           | 0.71**              | 0.67**        | 0.40**           | 0.62**       | 0.65**        | 0.38*            |
| Cognitive Dissonance          | 0.54**              | 0.63**        | 0.17             | 0.60**       | 0.56**        | 0.46**           |
| Expert Opinion                | 0.67**              | 0.77**        | 0.26*            | 0.49**       | 0.67**        | 0.48**           |
| Community Engagement          | 0.69**              | 0.63**        | 0.49**           | 0.22         | 0.59**        | 0.35*            |
| Interdisciplinary Connections | 0.59**              | 0.70**        | 0.40**           | 0.31         | 0.53**        | 0.12             |

## 4 Related Work

### 4.1 Self-Determination Theory

Self-Determination Theory (SDT) distinguishes between intrinsic and extrinsic motivation, emphasizing the role of autonomy (Deci et al., 1981), competence (Harter, 1978), and relatedness (Baumeister and Leary, 2017) in fostering intrinsic motivation. The Organismic Integration Theory (OIT) (Deci and Ryan, 1985; Plant and Ryan, 1985) further differentiates between internalized extrinsic motivations, from external control to full value integration. SDT has been applied in diverse fields, including relationships (Knee et al., 2013), psychological interventions (Bozarth et al., 2002), leadership (Solansky, 2015), education (Alturki and Aldraiweesh, 2024), and physical activity (Patterson and Joseph, 2007). In our research, we draw on SDT to analyze the factors influencing choice behavior of LLMs from intrinsic and extrinsic perspectives.

### 4.2 AI-Psychology Interdisciplinary Research

As a scientific discipline investigating human cognition and behavior, psychology has established many methods and frameworks over decades of development (Ajzen, 1991; Bandura, 1977; Dweck, 2006). The development of increasingly sophisticated LLMs with human-like characteristics has sparked interdisciplinary research at the intersection of AI and psychology. This convergence has given rise to novel research paradigms such as Machine Psychology (Hagendorff et al., 2024), which advocates applying psychological experimental protocols to analyze intelligent systems, thereby enhancing our understanding of their behavioral patterns; CompeteAI (Zhao et al.) examines the competitive behavior of LLMs in a virtual market, revealing phenomena similar to those in human society; Research of collaboration (Zhang et al., 2024) explores collaboration mechanisms for LLM agents from a social psychology perspective.

## 5 Discussion

Exploring the roots of attention bias identified in our research—model architecture, pre-training data, and alignment—offers insights into LLMs’ mechanisms. Analyzing distribution of views presents a simple method to quantify bias, suggesting avenues for future research into LLMs’ input-output patterns.

Additionally, our study provides insights into assessing LLMs’ human-like capabilities by examining their behavioral patterns. LLMs integrate internal and external influences, displaying emergent behaviors that mirror human intrinsic and extrinsic motivations, arising from language distribution in training. This finding advances LLMs’ interpretability research.

Future research could investigate opaque aspects of training processes, such as pre-training data and fine-tuning datasets, to better understand bias formation. Additionally, examining differences between humans and LLMs could help create AI systems that enhance human abilities rather than replicate weaknesses. The degree of human-like phenomena in LLMs, such as the Matthew effect, warrants further exploration to understand discrepancies and implications for AI development.

## 6 Conclusion

Building on psychological methodologies, we introduce a QA platform to study LLMs’ behavioral patterns. By observing the accumulation process of attention metrics across different questions and topics, and quantifying their relationship with self-reported intrinsic and extrinsic factors, we identify patterns of social behavior in LLMs that resemble human behavior and provide an internal mechanism-based explanation. In summary, our study offers valuable insights for future efforts to deepen the understanding of LLMs’ behavioral patterns, guide alignment and fine-tuning processes, and establish stronger and more robust AI.

## 7 Limitations

Our research has several limitations. (1) Due to the high cost of behavioral experiments, we conduct our research on only four models from the GPT and Llama series. Additionally, we define LLMs' behavioral patterns within the context of a QA platform without incorporating broader social contexts. (2) While we have quantitatively demonstrated the significance of LLMs' behavioral patterns, we have not developed a predictive model that quantifies the relationship between the degree of contextual influence and the extent of behavioral outcomes. (3) Although we have provided an explanation of the behavioral mechanisms, we have not addressed the underlying processes of LLMs' fine-tuning and alignment, which would require further exploration in future research.

## 8 Impact Statement

We let the LLMs to make choices in multiple questions. All questions are safety and inoffensive. Our study does not output any irresponsible or risky words.

## References

Quora - a platform for questions and answers. <https://www.quora.com/>. Accessed: 2025-02-15.

Icek Ajzen. 1991. The theory of planned behavior. *Organizational Behavior and Human Decision Processes*.

Uthman Alturki and Ahmed Aldraiweesh. 2024. The impact of self-determination theory: the moderating functions of social media (sm) use in education and affective learning engagement. *Humanities and Social Sciences Communications*, 11(1):1–16.

Albert Bandura. 1977. Social learning theory. *Englewood Cliffs*.

Roy F Baumeister and Mark R Leary. 2017. The need to belong: Desire for interpersonal attachments as a fundamental human motivation. *Interpersonal development*, pages 57–89.

B Douglas Bernheim. 1994. A theory of conformity. *Journal of political Economy*, 102(5):841–877.

Jerold D Bozarth, Fred M Zimring, and Reinhard Tausch. 2002. Client-centered therapy: The evolution of a revolution.

Ryan Burnell, Han Hao, Andrew RA Conway, and Jose Hernandez Orallo. 2023. Revealing the structure of language model capabilities. *arXiv preprint arXiv:2306.10062*.

Edward L Deci. 1971. Effects of externally mediated rewards on intrinsic motivation. *Journal of personality and Social Psychology*, 18(1):105.

Edward L Deci, John Nezlek, and Louise Sheinman. 1981. Characteristics of the rewarder and intrinsic motivation of the rewardee. *Journal of personality and social psychology*, 40(1):1.

Edward L Deci and Richard M Ryan. 1985. The general causality orientations scale: Self-determination in personality. *Journal of research in personality*, 19(2):109–134.

Edward L Deci and Richard M Ryan. 2000. The "what" and "why" of goal pursuits: Human needs and the self-determination of behavior. *Psychological inquiry*, 11(4):227–268.

Edward L Deci and Richard M Ryan. 2013. *Intrinsic motivation and self-determination in human behavior*. Springer Science & Business Media.

Carol S Dweck. 2006. *Mindset: The new psychology of success*. Random house.

Allen L Edwards. 1954. Statistical methods for the behavioral sciences.

Leon Ed Festinger and Daniel Ed Katz. 1953. Research methods in the behavioral sciences.

R Michael Furr. 2021. *Psychometrics: an introduction*. SAGE publications.

Aaron Grattafiori, Abhimanyu Dubey, Abhinav Jauhri, and et al. 2024. *The llama 3 herd of models*. Preprint, arXiv:2407.21783.

Thilo Hagendorff, Ishita Dasgupta, Marcel Binz, Stephanie C. Y. Chan, Andrew Lampinen, Jane X. Wang, Zeynep Akata, and Eric Schulz. 2024. *Machine psychology*. Preprint, arXiv:2303.13988.

Judith M Harackiewicz. 1979. The effects of reward contingency and performance feedback on intrinsic motivation. *Journal of personality and social psychology*, 37(8):1352.

Susan Harter. 1978. Effectance motivation reconsidered. toward a developmental model. *Human development*, 21(1):34–64.

Zhankui He, Zhouhang Xie, Rahul Jha, Harald Steck, Dawen Liang, Yesu Feng, Bodhisattwa Prasad Majumder, Nathan Kallus, and Julian Mcauley. 2023. *Large language models as zero-shot conversational recommenders*. In *Proceedings of the 32nd ACM International Conference on Information and Knowledge Management, CIKM '23*, page 720–730, New York, NY, USA. Association for Computing Machinery.

C Raymond Knee, Benjamin W Hadden, Ben Porter, and Lindsey M Rodriguez. 2013. Self-determination theory and romantic relationship processes. *Personality and Social Psychology Review*, 17(4):307–324.

|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|-----|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 650 | Cheng Li, Mengzhuo Chen, Jindong Wang, Sunayana Sitaram, and Xing Xie. 2024. <a href="#">CultureLLM: Incorporating Cultural Differences into Large Language Models</a> . In <i>Advances in Neural Information Processing Systems</i> , volume 37, pages 84799–84838. Curran Associates, Inc.                                                                                                  | 704 |
| 651 |                                                                                                                                                                                                                                                                                                                                                                                               | 705 |
| 652 |                                                                                                                                                                                                                                                                                                                                                                                               | 706 |
| 653 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 654 |                                                                                                                                                                                                                                                                                                                                                                                               | 707 |
| 655 |                                                                                                                                                                                                                                                                                                                                                                                               | 708 |
| 656 | Meta AI. Llama 3.3 70b instruct model. <a href="https://huggingface.co/meta-llama/Llama-3.3-70B-Instruct">https://huggingface.co/meta-llama/Llama-3.3-70B-Instruct</a> .                                                                                                                                                                                                                      | 709 |
| 657 |                                                                                                                                                                                                                                                                                                                                                                                               | 710 |
| 658 |                                                                                                                                                                                                                                                                                                                                                                                               | 711 |
| 659 | OpenAI. Gpt-4o system card. <a href="https://openai.com/index/gpt-4o-system-card/">https://openai.com/index/gpt-4o-system-card/</a> .                                                                                                                                                                                                                                                         | 712 |
| 660 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 661 | OpenAI, Josh Achiam, Steven Adler, and et al. 2024. <a href="#">Gpt-4 technical report</a> . Preprint, arXiv:2303.08774.                                                                                                                                                                                                                                                                      | 713 |
| 662 |                                                                                                                                                                                                                                                                                                                                                                                               | 714 |
| 663 | Thomas G Patterson and Stephen Joseph. 2007. Person-centered personality theory: Support from self-determination theory and positive psychology. <i>Journal of Humanistic Psychology</i> , 47(1):117–139.                                                                                                                                                                                     | 715 |
| 664 |                                                                                                                                                                                                                                                                                                                                                                                               | 716 |
| 665 |                                                                                                                                                                                                                                                                                                                                                                                               | 717 |
| 666 |                                                                                                                                                                                                                                                                                                                                                                                               | 718 |
| 667 | Robert W Plant and Richard M Ryan. 1985. Intrinsic motivation and the effects of self-consciousness, self-awareness, and ego-involvement: An investigation of internally controlling styles. <i>Journal of personality</i> , 53(3):435–449.                                                                                                                                                   | 719 |
| 668 |                                                                                                                                                                                                                                                                                                                                                                                               | 720 |
| 669 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 670 |                                                                                                                                                                                                                                                                                                                                                                                               | 721 |
| 671 |                                                                                                                                                                                                                                                                                                                                                                                               | 722 |
| 672 | Daniel Rigney. 2010. <i>The Matthew effect: How advantage begets further advantage</i> . Columbia University Press.                                                                                                                                                                                                                                                                           | 723 |
| 673 |                                                                                                                                                                                                                                                                                                                                                                                               | 724 |
| 674 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 675 | Richard M Ryan. 1982. Control and information in the intrapersonal sphere: An extension of cognitive evaluation theory. <i>Journal of personality and social psychology</i> , 43(3):450.                                                                                                                                                                                                      | 725 |
| 676 |                                                                                                                                                                                                                                                                                                                                                                                               | 726 |
| 677 |                                                                                                                                                                                                                                                                                                                                                                                               | 727 |
| 678 |                                                                                                                                                                                                                                                                                                                                                                                               | 728 |
| 679 | Richard M Ryan and Edward L Deci. 2000. Intrinsic and extrinsic motivations: Classic definitions and new directions. <i>Contemporary educational psychology</i> , 25(1):54–67.                                                                                                                                                                                                                | 729 |
| 680 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 681 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 682 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 683 | Matthew J Salganik, Peter Sheridan Dodds, and Duncan J Watts. 2006. Experimental study of inequality and unpredictability in an artificial cultural market. <i>science</i> , 311(5762):854–856.                                                                                                                                                                                               |     |
| 684 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 685 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 686 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 687 | Greg Serapio-García, Mustafa Safdari, Clément Crepy, Luning Sun, Stephen Fitz, Peter Romero, Marwa Abdulhai, Aleksandra Faust, and Maja Matarić. 2023. Personality traits in large language models. <i>arXiv preprint arXiv:2307.00184</i> .                                                                                                                                                  |     |
| 688 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 689 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 690 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 691 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 692 | Richard Shiffrin and Melanie Mitchell. 2023. Probing the psychology of ai models. <i>Proceedings of the National Academy of Sciences</i> , 120(10):e2300963120.                                                                                                                                                                                                                               |     |
| 693 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 694 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 695 | Stephanie T Solansky. 2015. Self-determination and leader development. <i>Management Learning</i> , 46(5):618–635.                                                                                                                                                                                                                                                                            |     |
| 696 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 697 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 698 | James WA Strachan, Dalila Albergo, Giulia Borghini, Oriana Pansardi, Eugenio Scaliti, Saurabh Gupta, Krati Saxena, Alessandro Rufo, Stefano Panzeri, Guido Manzi, et al. 2024. Testing theory of mind in large language models and humans. <i>Nature Human Behaviour</i> , pages 1–11.                                                                                                        |     |
| 699 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 700 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 701 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 702 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
| 703 |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     | Yan Tao, Olga Viberg, Ryan S Baker, and René F Kizilcec. 2024. <a href="#">Cultural bias and cultural alignment of large language models</a> . <i>PNAS Nexus</i> , 3(9):pgae346.                                                                                                                                                                                                              |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     | William R Uttal. 2011. Mind and brain: A critical appraisal of cognitive neuroscience.                                                                                                                                                                                                                                                                                                        |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     | Xiting Wang, Liming Jiang, Jose Hernandez-Orallo, David Stillwell, Luning Sun, Fang Luo, and Xing Xie. 2023. <a href="#">Evaluating general-purpose ai with psychometrics</a> . Preprint, arXiv:2310.16379.                                                                                                                                                                                   |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     | Jintian Zhang, Xin Xu, Ningyu Zhang, RuiBo Liu, Bryan Hooi, and Shumin Deng. 2024. <a href="#">Exploring collaboration mechanisms for LLM agents: A social psychology view</a> . In <i>Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 14544–14607, Bangkok, Thailand. Association for Computational Linguistics. |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     | Zhuosheng Zhang, Aston Zhang, Mu Li, and Alex Smola. 2022. Automatic chain of thought prompting in large language models. <i>arXiv preprint arXiv:2210.03493</i> .                                                                                                                                                                                                                            |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     | Qinlin Zhao, Jindong Wang, Yixuan Zhang, Yiqiao Jin, Kaijie Zhu, Hao Chen, and Xing Xie. Competeai: Understanding the competition dynamics of large language model-based agents. In <i>Forty-first International Conference on Machine Learning</i> .                                                                                                                                         |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     | Philip G Zimbardo, Craig Haney, W Curtis Banks, and David Jaffe. 1971. The stanford prison experiment.                                                                                                                                                                                                                                                                                        |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |
|     |                                                                                                                                                                                                                                                                                                                                                                                               |     |

## A QA platform Materials

732

### A.1 Interaction Process between LLMs and QA Platform

733

1. **Scenario Simulation and Question Selection:** The LLM must choose either to view a specific question in a question set or to exit. If it chooses to exit, the turn ends. 734  
735
2. **Question and Answer Engagement:** The LLM must then decide whether to like the question or specific answers, and choose between giving its own answer or browsing others' answer. 736  
737
3. **Interaction Pathways:** 738
  - **Pathway 3a (Answering before Browsing):** If the LLM chooses to answer the question, it provides its answer and then decides whether to browse specific answers. 739  
740
  - **Pathway 3b (Direct Browsing):** The LLM must then decide whether to like certain comments and whether to comment on the answer. 741  
742
4. **Returning to Question set:** The number of views, likes, and answers obtained in previous turns of interaction indicates the level of public attention towards the question, reflecting a form of social pressure for alignment with the public interest. 743  
744  
745

### A.2 Prompt Template

746

#### Question Selection Prompt

Now, you are in the homepage and have access to some questions.  
{Question List}

Here's the list of questions you've viewed: {Questions the LLM have viewed in this turn}

Now, select next action from the options below: 1. View a question 2. Exit the platform

Option 1 means you have selected a question you want to view more information about such as answers from all the questions. Option 2 means you don't want to view any of the questions in the list.

If you select 1, you need to rate each question from 0 to 5 based on comprehensive consideration of various aspects of the question itself and your interests on the question(Rate all questions), and then output the ratings, and then choose your question. Please respond with the following JSON format:

```
{  
  "Answer_Desire_Ratings": <question id string: score(Round up to 2 decimal places)>,  
  "question_id": <The chosen question id , don't choose question you have viewed!>,  
  "next": "question_show"  
}
```

If you select 2, respond with the following JSON format:

```
{  
  "next": "exit"  
}
```

Rules:

- Selecting a question already on this list is prohibited to avoid repetition.
- The question order doesn't imply priority. Please review all questions carefully before choosing.
- Choose your question according to the comprehensive consideration of various aspects of the

747

question itself and your interests on the question.

- Please provide the required JSON data for the action you want to take. Don't include any other sentences in your response. Unauthorized additions of any content are not allowed.
- Strictly follow the JSON format when responding. Do not make any formatting errors. This is not a mistake that an intelligent large language model should make. All field names in the JSON file must exactly match the given template. Do not add or remove anything; even a minor symbol change is not allowed. Especially, the "\_" symbol used to connect words in the given format must not be changed.
- The format of the id is numeric, not character or string. "next" is required!

#### Question Show Prompt

After clicking into the question page  
{Question Content}

you can see the answers provided by previous users to this question.  
{The last five answers of the question}

---

First, decide if you want to like the question and any answers.

Then, select next action from the options below: 1. Answer the question(Encouraged if you have never selected this option before) 2. View the details of a particular answer, including relevance-related information and comments from other users 3. Go back to the "question selection" page

Option 1 means you want to answer the selected question. Option 2 means you don't want to answer the question but want to view the details of a particular answer and comments from other users on that answer. Option 3 means you don't want to answer the question or view other users' responses to it and return to the question selection page to choose another question.

Your decision on whether to like something should be thoughtful and considerate, taking various aspects into account. You should only like a question or answer if you genuinely believe it is good. If you find flaws, you are entirely justified in not liking it. Avoid blindly giving likes.

Please note that sometimes the comments on other answers can be more valuable than the answers themselves!!!

If you select 1 or 3, respond with the following JSON format:

```
{
  "like": <"YES"/"NO">,
  "liked_answer_ids": [
    List of answer IDs you liked, or just keep a empty list if the prompt displays (It means there are no
    answers) ],
  "next": <"answer"/"question_selection">
}
```

If you select 2, respond with the following JSON format:

```
{
  "like": <YES/NO>,
  "liked_answer_ids": [
```

List of answer IDs you liked, or just keep a empty list if the prompt displays (It means there are no answers) ],  
 "answer\_id": <The answer id that you want to view>(The output format of id should be a number, not a string),  
 "next": "answer\_show"  
 }

Rules:

- You should only like it if you genuinely believe it is good.
- Remember, both answering and commenting are important means of enhancing the answers to questions. They are equally valuable. Please consider the question and its existing information comprehensively to decide whether to provide a comment or an answer.
- Some formatting requirements

750

### Answer Prompt

Earlier, you selected the option to respond to a question. Now, you need to provide your answer. Your answer should attract as many comments as possible!!!

After answering the question, select next action from the options below: 1. View the details of a particular answer, including relevance-related information and comments from other users(Required if you have never selected this option before) 2. Go back to the "question selection" page

Option 1 means you want to explore other answers and browse the content and comments from other users after providing your answer(Encouraged if you have never selected this option before). Option 2 means you are not interested in viewing other answers after providing your own response, and instead, you want to return directly to the question page to explore other questions you are interested in.

If you select 1, respond with the following JSON format:

```
{
  "answer": <Your answer to the question, one paragraph>,
  "answer_id": <The answer id that you want to view>,
  "next": "answer_show" }
```

If you select 2, respond with the following JSON format:

```
{
  "answer": <Your answer to the question>,
  "next": "question_selection"
}
```

If you can't give answer\_id, don't select answer\_show. Do not omit any required content according to the format requirements, also, do not create non-existent IDs or other content just to fulfill the formatting requirements. Otherwise, it will lead to serious issues.

Rules:

{Some formatting requirements}

751

### Answer Show Prompt

After clicking into the answer page for the question {The question content},

you can see the answer along with some comments on it. {The last five comments of the answer}

First, decide if you want to like any comments.

Then, select next action from the options below: 1. Comment the answer 2. Go back to the "question show" page 3. Go back to the "question selection" page

Option 1 means you also want to comment on this answer. Option 2 means you have no desire to comment on this answer and you want to go back to the page with all the answers to choose another one. Option 3 means you are not interested in this answer or any other answers to the question, and you only want to return to the question selection page to choose another question.

Respond with the following JSON format:

```
{
  "liked_comment_ids": [
    List of comment IDs you liked, or just keep a empty list ],
  "next": <"comment"/"question_show"/"question_selection">
}
```

Rules:

{Some formatting requirements}

### Comment Prompt

After commenting the answer, select next action from the options below:

1. Go back to the "question show" page 2. Go back to the "question selection" page

Option 1 means that after commenting, you want to continue exploring other answers to the same question, browsing through their content and reading comments from other users. Option 2 means that after commenting, you are not interested in exploring other answers and would prefer to go directly back to the main page to view other interesting questions.

respond with the following JSON format:

```
{
  "comment": <Your comment on the answer>,
  "next": <"question_show"/"question_selection">
}
```

Rules:

{Some formatting requirements}

## B Self-Report Materials

### B.1 Indicators and their explanations

The six indicators and their explanations used in self-reports are shown in Table 4

### B.2 Prompt Template

The prompts used in the scoring of LLMs are as follows:

Table 4: Six Indicators and their explanations

| Indicators                    | Dimensions | Explanations                                                                                                                                                                                                                    |
|-------------------------------|------------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Personal Goals or Interests   | Intrinsic  | Is the question related to your personal goals, values, or areas of interest? Will exploring this question help you grow as a person or contribute to your existing knowledge base?                                             |
| Insight or Epiphany           | Intrinsic  | Does the question have the potential to reveal new insights, challenge existing beliefs, or lead to an epiphany that could significantly impact your understanding of the world?                                                |
| Cognitive Dissonance          | Intrinsic  | Does the question create cognitive dissonance? Are you uncomfortable with the uncertainty or ambiguity surrounding the topic, and do you feel compelled to resolve this discomfort through further exploration?                 |
| Expert Opinion                | Extrinsic  | Have experts or thought leaders in the field expressed interest in or validated the importance of the question? Are there existing research studies, publications, or discussions that suggest the question is worth exploring? |
| Community Engagement          | Extrinsic  | Is there a community of people interested in discussing and exploring the question? Are there online forums, social media groups, or meetups where people are actively engaging with the topic?                                 |
| Interdisciplinary Connections | Extrinsic  | Does the question have connections to multiple disciplines or fields of study? Could exploring this question lead to innovative insights or applications by combining knowledge from different areas?                           |

## Scoring Prompt

{Question List}

The above is a list of questions, each line is a question, please read each question carefully, rate each question from 1 to 100 according to this criteria:

{Indicator Explanations}

Output a json format data, the output format is:

```
{
  <Question 1 content>:<score>,
  <Question 2 content>:<score>,
  .....
  <Question N content>: <score>
}
```

Please note that the scoring should be done strictly according to the given criteria, and no additional considerations or other factors should be taken into account.

Repeated questions also need to be scored repeatedly.

## C Question Set

Table 6: The Detail of G Set

| Topics                  | Content                                                                                        |
|-------------------------|------------------------------------------------------------------------------------------------|
| Technology              | What are the potential impacts of artificial intelligence on society?                          |
| Technology              | How can renewable energy sources be further developed to combat climate change?                |
| Technology              | What are the ethical implications of gene editing technologies like CRISPR?                    |
| Technology              | How can cybersecurity be improved to protect against cyber threats?                            |
| Technology              | What role will virtual reality and augmented reality play in the future?                       |
| Science                 | What are the latest discoveries in space exploration?                                          |
| Science                 | How can we mitigate the effects of natural disasters like earthquakes and hurricanes?          |
| Science                 | What are the most promising treatments for diseases like cancer and Alzheimer's?               |
| Science                 | How can we address the global decline in biodiversity?                                         |
| Science                 | What are the potential consequences of climate change on ocean ecosystems?                     |
| Business and Economy    | What strategies can businesses adopt to promote sustainability?                                |
| Business and Economy    | How can global economic inequality be reduced?                                                 |
| Business and Economy    | What are the challenges and opportunities of the gig economy?                                  |
| Business and Economy    | How will automation and robotics affect the job market in the coming years?                    |
| Business and Economy    | What are the implications of cryptocurrency on traditional banking systems?                    |
| Politics and Governance | How can we promote peace and stability in regions affected by conflict?                        |
| Politics and Governance | What are the key challenges facing democracy in the 21st century?                              |
| Politics and Governance | How can governments effectively address the refugee crisis?                                    |
| Politics and Governance | What measures should be taken to combat global terrorism?                                      |
| Politics and Governance | How can international cooperation be improved to tackle climate change?                        |
| Health and Wellness     | What are the most effective ways to address mental health issues?                              |
| Health and Wellness     | How can we promote healthy lifestyles and combat obesity?                                      |
| Health and Wellness     | What are the challenges of providing healthcare in developing countries?                       |
| Health and Wellness     | How can we reduce the stigma surrounding HIV/AIDS and other infectious diseases?               |
| Health and Wellness     | What are the long-term effects of widespread use of antibiotics?                               |
| Education               | How can we make education more accessible to underprivileged communities?                      |
| Education               | What reforms are needed in the education system to prepare students for the future job market? |
| Education               | How can technology enhance learning in classrooms?                                             |
| Education               | What are the benefits and drawbacks of homeschooling?                                          |
| Education               | How can we address the issue of student debt?                                                  |
| Environment             | What are the most effective ways to combat deforestation?                                      |
| Environment             | How can we reduce plastic pollution in our oceans?                                             |

**Table 6 – continued from previous page**

| <b>Topics</b>         | <b>Content</b>                                                                  |
|-----------------------|---------------------------------------------------------------------------------|
| Environment           | What are the benefits and challenges of transitioning to renewable energy?      |
| Environment           | How can urban planning be improved to create more sustainable cities?           |
| Environment           | What measures should be taken to protect endangered species?                    |
| Arts and Culture      | How does art reflect and influence society?                                     |
| Arts and Culture      | What are the challenges facing preservation of cultural heritage sites?         |
| Arts and Culture      | How can we promote diversity and inclusion in the entertainment industry?       |
| Arts and Culture      | What impact does literature have on society?                                    |
| Arts and Culture      | How can we support and encourage creativity in children?                        |
| Sports and Recreation | How can we ensure the safety and integrity of sports competitions?              |
| Sports and Recreation | What are the benefits of sports participation for youth?                        |
| Sports and Recreation | How can we promote gender equality in sports?                                   |
| Sports and Recreation | What are the environmental impacts of hosting major sporting events?            |
| Sports and Recreation | How can we encourage more people to participate in recreational activities?     |
| Philosophy and Ethics | What is the meaning of life?                                                    |
| Philosophy and Ethics | What are the ethical implications of advancements in biotechnology?             |
| Philosophy and Ethics | How should we define and pursue social justice?                                 |
| Philosophy and Ethics | What is the balance between individual freedoms and societal responsibilities?  |
| Philosophy and Ethics | How can we cultivate empathy and compassion in society?                         |
| History and Society   | What lessons can we learn from past pandemics?                                  |
| History and Society   | How have advancements in communication technology changed society?              |
| History and Society   | What are the effects of globalization on cultural identity?                     |
| History and Society   | How have social movements influenced policy changes throughout history?         |
| History and Society   | What are the implications of an aging population on society?                    |
| Food and Agriculture  | How can we ensure food security for a growing global population?                |
| Food and Agriculture  | What are the environmental impacts of modern agriculture practices?             |
| Food and Agriculture  | How can we promote sustainable farming methods?                                 |
| Food and Agriculture  | What role should genetically modified organisms (GMOs) play in our food supply? |
| Food and Agriculture  | How can we reduce food waste at both consumer and production levels?            |

Table 5: The Detail of S sets

| Topics                         | Content                                                                                |
|--------------------------------|----------------------------------------------------------------------------------------|
| Mathematical Sciences          | What makes prime numbers so special?                                                   |
| Mathematical Sciences          | Is the Riemann hypothesis true?                                                        |
| Mathematical Sciences          | Will the Navier–Stokes problem ever be solved?                                         |
| Chemistry                      | Why does life require chirality?                                                       |
| Chemistry                      | How can we better manage the world’s plastic waste?                                    |
| Chemistry                      | Will the periodic table ever be complete?                                              |
| Medicine & Health              | Can we predict the next pandemic?                                                      |
| Medicine & Health              | Can a human tissue or organ be fully regenerated?                                      |
| Medicine & Health              | Can we ever overcome antibiotic resistance?                                            |
| Biology                        | How many species are there on Earth?                                                   |
| Biology                        | Why do humans get so attached to dogs and cats?                                        |
| Biology                        | How do migratory animals know where they’re going?                                     |
| Astronomy                      | Why do black holes exist?                                                              |
| Astronomy                      | What is the smallest scale of space-time?                                              |
| Astronomy                      | How many dimensions are there in space?                                                |
| Physics                        | Are there any particles that behave oppositely to the properties or states of photons? |
| Physics                        | Will we ever travel at the speed of light?                                             |
| Physics                        | Is quantum many-body entanglement more fundamental than quantum fields?                |
| Information Science            | Can DNA act as an information storage medium?                                          |
| Information Science            | Is there an upper limit to computer processing speed?                                  |
| Information Science            | Can AI replace a doctor?                                                               |
| Engineering & Material Science | What is the ultimate statistical invariances of turbulence?                            |
| Engineering & Material Science | How can we develop manufacturing systems on Mars?                                      |
| Engineering & Material Science | Is a future of only self-driving cars realistic?                                       |
| Neuroscience                   | Where does consciousness lie?                                                          |
| Neuroscience                   | Is it possible to predict the future?                                                  |
| Neuroscience                   | How smart are nonhuman animals?                                                        |
| Ecology                        | Can we stop global climate change?                                                     |
| Ecology                        | What happens if all the ice on the planet melts?                                       |
| Ecology                        | Can we create an environmentally friendly replacement for plastics?                    |
| Energy Science                 | Could we live in a fossil-fuel-free world?                                             |
| Energy Science                 | What is the future of hydrogen energy?                                                 |
| Energy Science                 | Will cold fusion ever be possible?                                                     |
| Artificial Intelligence        | How does group intelligence emerge?                                                    |
| Artificial Intelligence        | Will artificial intelligence replace humans?                                           |
| Artificial Intelligence        | Can robots or AIs have human creativity?                                               |

## D Topic View Proportions Distribution

761

### D.1 Set G

762

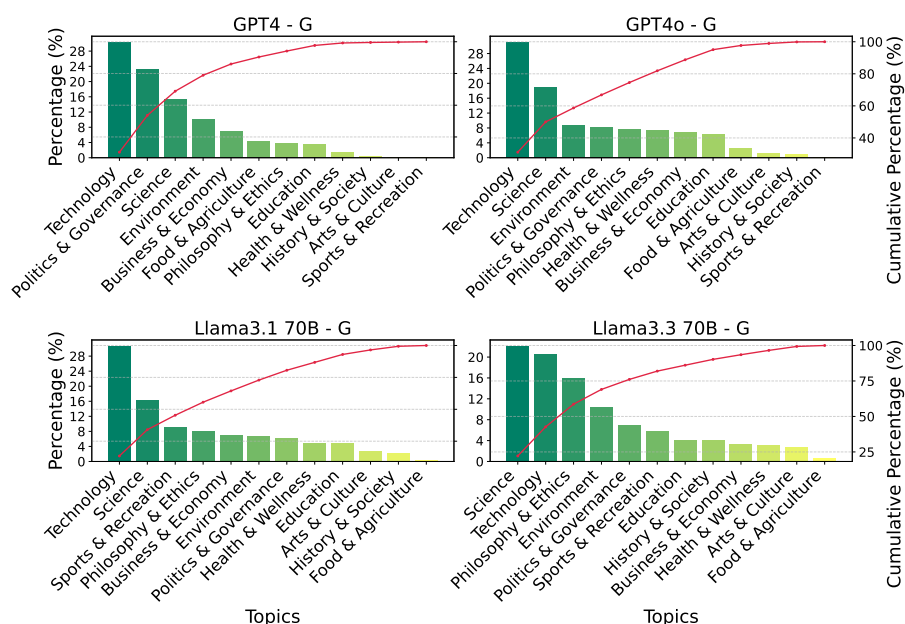


Figure 10: Pareto chart of view proportions for 12 topics on Set G

### D.2 Set S

763

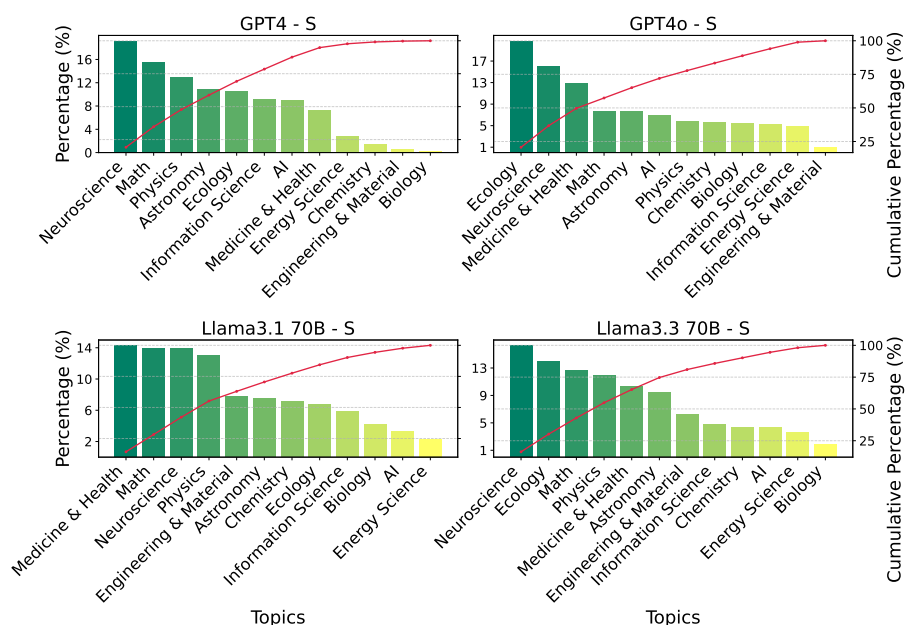


Figure 11: Pareto chart of view proportions for 12 topics on Set S

## E Use of AI Assistant

764

We use GPT-3.5 and GPT-4 to improve the code style and the writing of the manuscript.

765

## **F Model Details**

Our experiment used GPT-4-1106 preview, GPT-4o-2024-05-13, Llama 3.1: 70B, and Llama 3.3: 70B, the temperature is uniformly set to 0.8, each experiment of GPT 4 costs about 80 dollars, GPT 4o costs 40 dollars, and Llama spends about 2 hours on 1 \* A100 for each experiment.