

Understanding Fortunetelling with Large Language Models in China: User Practices, Perceptions, and Impacts on Beliefs and Decisions

Xueer Lin¹, Chenyu Li¹, Shuai Ma³, Yuhao Lyu¹, Zhenhui Peng^{1*}

¹Sun Yat-sen University, Zhuhai, China

²Aalto University, Helsinki, Finland

linxer6@mail2.sysu.edu.cn, lichy235@mail2.sysu.edu.cn, shuai.ma@connect.ust.hk, lvyh26@mail2.sysu.edu.cn, pengzh29@mail.sysu.edu.cn

Abstract

Fortunetelling is a cultural practice for navigating uncertainty, often associated with people’s beliefs and decisions. Fortunetelling with recent large language models (LLMs) introduces new opportunities and risks. This paper conducts qualitative studies to understand users’ practices, perceptions, and impacts of LLM fortunetelling in China. We first analyze 1,045 posts on Chinese social media, yielding a comprehensive taxonomy of the diverse foretold topics (*e.g.*, career, romance), emotion reactions (*e.g.*, surprise, worry), and perceived credibility (*e.g.*, doubt, trust) of LLM fortunetelling. Then, we conduct interviews with 20 users of LLM fortunetelling. The findings indicate that users treat LLM fortunetelling as a tool less for accurate prediction but more for emotional support. While the fortunetelling results rarely change users’ initial beliefs or decisions, they are associated with subtle mindset shifts, with some users reporting small behavioral adjustments. We discuss implications for gaining benefits from LLM fortunetelling.

Introduction

Fortunetelling, which attempts to predict or reveal information about a person’s future, character, or fate through symbolic systems, rituals, or intuition, has long been practiced in many cultures (Smith 2019). Specifically, in Chinese culture focused in this paper, fortunetelling is traditionally facilitated by human masters, conducted via metaphysical systems like Bazi and Ziwei Doushu (Purple Star Astrology), and used for important decisions such as marriage matching and naming children. Recent advances in large language models (LLMs) have boost a new accessible form of fortunetelling, *i.e.*, *LLM fortunetelling*, actively discussed on social media platforms. People increasingly turn to LLMs to seek predictions about major life events (*e.g.*, romantic prospects, academic success) and share these experiences on social media. For instance, on RedNote — a popular Chinese platform — users often prompt LLMs to act as fortunetelling masters by providing personal details such as birth date, time, and zodiac, and then discuss the predictions with others (Figure 1).

This phenomenon introduces both opportunities and risks. On the one hand, LLM fortunetelling can provide emotional support, stimulate self-reflection, and guide daily decision-making (Tan, Van Prooijen, and van Lange 2022). On the other hand, prior ICWSM work has already raised concerns about LLMs in sensitive advice settings (*e.g.*, relationship advice) (Hou, Leach, and Huang 2024). LLM fortunetelling inherits challenges of LLMs such as hallucinations (Meier 2022) and may encourage over-reliance, potentially leading to poor financial, academic, or career-related decisions (Grall-Bronnec et al. 2015). This “double-edged sword” calls for a deeper understanding of how people experience, perceive, and respond to LLM fortunetelling, so that its benefits can be harnessed while mitigating its risks.

However, the use of LLMs for fortunetelling — particularly how users engage with its predictions — remains largely unexplored. First, although LLM fortunetelling is often treated as a playful use case, prior research suggests that casual interactions with LLMs can acquire personal meaning and shape users’ beliefs (Danry et al. 2025). Second, LLM fortunetelling is increasingly invoked for decision-making support. Yet, existing studies on AI-assisted decision-making have largely centered on expert or objective domains such as clinical diagnosis (Tsai et al. 2021) or financial forecasting (Li et al. 2025). Unlike these contexts, fortunetelling concerns users’ own lives and choices, making AI-generated responses personally meaningful. Investigating these practices can therefore deepen our understanding of AI-mediated decision-making and belief formation in self-referential contexts.

To address these gaps, we investigate the following research questions in the Chinese context: **RQ1**: What are the user practices of LLM fortunetelling regarding the foretold topics, prompts, and motivations? **RQ2**: What are the user perceptions of LLM fortunetelling? **RQ3**: How would the users’ beliefs and decisions on the foretold topics be affected by LLM fortunetelling?

To answer these questions, we first analyzed 1,045 posts about fortunetelling with LLMs from RedNote and Weibo — two representative social media platforms in China, providing insights from a large and diverse user base to address RQ1 and RQ2. Then, we conducted interviews with 20 users of LLM fortunetelling to gain in-depth qualitative answers to each RQ. These two qualitative studies yield a compre-

*Corresponding author.



Figure 1: Three topics of posts. Post content is translated from Chinese to English using the embedded translator of the browser.

hensive taxonomy of user practices and perceptions of LLM fortunetelling. We found that people adopt well-structured prompts to consult LLMs on both traditional topics like Bazi (general fortune) and issues in modern society like pursuing postgraduate studies and locating lost property. As for user perceptions of and responses to LLM fortunetelling, our findings suggest the interplay of motivated reasoning, confirmation bias, and the Barnum effect, where users tend to interpret vague or personalized outputs in ways that reinforce their prior beliefs, expectations, and emotional needs (Nickerson 1998; Kunda 1990; Huizi, Yiwei et al. 2022). Most users were positive towards LLM fortunetelling, feeling surprised when predictions coincidentally matched reality and joking about the implausible results. Besides, participants in the interview study reported that they turned to LLM fortunetelling for both decision support and emotional comfort. When interacting with the LLM, they proactively managed risks and filtered the fortunetelling results.

In summary, this work contributes to a comprehensive understanding of LLM fortunetelling’s practices, user perceptions, and impacts on users’ beliefs and decisions. The findings reveal motivated reasoning, confirmation bias, and the Barnum effect in LLM fortunetelling, suggesting a new perspective on the role of LLM, via fortunetelling activities, in emotion regulation and decision support. Finally, this work highlights the potential implications of LLM fortunetelling on social media, where LLM-generated narratives shape how uncertainty is publicly expressed and managed, with subtle consequences over time.

Related Work

Fortunetelling as a Cultural Practice

Fortunetelling exists in diverse cultural contexts, from Eastern traditions like Bazi and Ziwei Doushu to Western customs like tarot cards and astrology. For example, in China, people consult Bazi not only to form beliefs about their life trajectory and future potential, but also to guide concrete decisions such as choosing an auspicious date for marriage. Psychological and consumer behavior research suggests that positive fortunetelling experiences can enhance subjective well-being and increase risk-taking (Tan, Van Prooijen, and van Lange 2022), while dependence on fortunetelling may negatively impact decision-making and mental health (Skryabin 2020).

The rise of large language models (LLMs) has made AI-mediated fortunetelling more accessible. Related research has explored models optimized for divination-like outputs (Cheng et al. 2025), while human-computer interaction researchers have developed multimodal conversational agents inspired by shamanic symbolism to study symbolic AI interaction (Cho et al. 2025). Recent studies on digital spirituality further show that social media enables the creation and circulation of symbolic meaning. For example, astrology content on social media shapes users’ self-concepts and relationships, while platforms such as TikTok have been described as “sacred technologies” that reinforce spiritual narratives (Zhang and Wang 2025). These related studies collectively indicate the trends of using technology to facilitate traditional fortunetelling practices. However, there is a lack of systematic understanding of users’ experiences with the advanced LLM fortunetelling. Our research fills this gap by providing insights into how users turn to LLMs to foretell everyday issues, how they feel about LLM fortunetelling, and how their beliefs and decisions are affected.

LLM in Playful Contexts, for Emotional Support, and Its Impact on Human Beliefs

Previous research has integrated LLMs into informal everyday contexts, indicating that even seemingly playful interactions with LLMs can have unexpected meanings and consequences. For example, research on BleacherBot (Kim et al. 2025) shows that AI can serve as a companion while watching sports. Because it is not a real person, users feel more comfortable without the usual social pressures of watching games with others. These findings collectively highlight the characteristics of LLM use: interactions may appear playful, but can subtly influence users’ emotions, reflections, and decisions. More recent work further illustrates the potential of LLM to provide emotional support and affect user beliefs. For example, research on ChatLab (Zheng et al. 2025) shows how users construct personalized LLM-powered chatbot personas to confront stressors, foster emotional reliance, and engage in self-reflection. Empirical evidence further illustrates how LLM behaviors can shape belief formation. The research conducted by Danry et al. (2025) finds that deceptive AI-generated explanations are more persuasive than accurate ones and could significantly increase belief in false news while eroding belief in true

news. Our work extends this understanding by investigating how seemingly playful interactions shape users’ emotions and beliefs in the context of LLM fortunetelling.

AI-Powered Decision-Making Support

A growing body of research has explored AI-powered decision-making, mostly in domains where decisions are made by experts, such as clinical diagnosis (Tsai et al. 2021) and income forecasting (Li et al. 2025), where professionals use AI systems to make decisions on behalf of others. These studies mainly focus on model performance, outcome quality, and issues of fairness and accountability. Recently, researchers have begun to examine AI-supported decision-making initiated by individuals. However, much of this research relies on low-stake tasks defined by the researchers. Common examples include choosing an exercise plan (Salimzadeh, He, and Gadiraju 2024) or planning a trip (Bućinca et al. 2025). These scenarios are often used to study trust, framing, or interpretation, but they often lack emotional depth or lasting impact. A relevant example is research on Quark Gaokao¹, a widely used AI tool for college admissions planning in China. Despite the importance of the decisions, research has found that parents often dominate the process of choosing universities, and students, while directly affected, have limited involvement (Chen et al. 2025). Furthermore, previous AI-powered decision-making support tools usually rely on ground-truth data or models trained on it.

In contrast, LLM fortunetelling does not rely on objective data sources for decision recommendations. Instead, it involves self-directed decisions that carry emotional and social significance. Rather than completing researcher-designed tasks, users engage with LLMs to explore uncertainty, seek guidance, and construct personal meaning. Our work contributes to understanding whether and how users’ decisions and behaviors are affected in the context of LLM fortunetelling.

Descriptive Content Analysis of Social Media Data about Fortunetelling with LLMs

We used two qualitative studies to understand fortunetelling with LLMs. Figure 2 shows our research flow. This section presents descriptive content analysis of posts in social media to gain a comprehensive understanding, while the next section presents an interview study to gain in-depth qualitative insights. Previous research has leveraged social media data to understand users’ practices and opinions regarding LLM usage in specific domains like education (Tlili et al. 2023). Besides, prior ICWSM work has examined how social media discourse reflects and shapes user behavior and social phenomena, including issues such as online incivility and youth well-being (Kim et al. 2024). This line of research underscores social media as a context in which individual experiences and expressions can become publicly visible and socially consequential. Following that, we collected data,

¹https://vt.quark.cn/blm/pc-gaokao-1089/index?uc_param_str=dnntnwvpeffrgibjbrsvdsdicheiniutskp&x_render_type=csr&ch=pcquark@homepage_gaokao_seoredirect&entry=p_redirect

conducted open coding, and trained classifiers, producing a comprehensive taxonomy of users’ practices (RQ1) and perceptions (RQ2) of LLM fortunetelling.

Data Collection

We collected data from Rednote (or Xiaohongshu) and Weibo, two major Chinese social media platforms. We selected these platforms because they form a complementary and widely adopted ecosystem, enabling us to capture diverse user practices surrounding LLM fortunetelling while accounting for platform-level differences through stratified analysis (see Table 1 and Table 2). As of the second quarter of 2025, Rednote hosts about 300 million monthly active users and over 100 million contributors², primarily young users sharing lifestyle content. In comparison, Weibo has 588 million monthly active users³ and mainly supports public discussion, trending topics, and rapid information sharing. To search posts relevant to fortunetelling with LLMs, we used combined keywords about mainstream LLMs (*e.g.*, “DeepSeek”, “GPT”, “Gemini”, “Grok”) and fortunetelling (*e.g.*, “fortunetelling”, “metaphysics”, “decision”, “choice”). From Rednote, we used the Spider_XHS crawler⁴ to retrieve the top 200 posts for each combined search keyword. After removing duplicated posts, we initially had 1,157 posts from Rednote. From Weibo, we manually collected 230 text-image posts using the combined keywords in the Weibo search engine. To ensure comparability across platforms, we applied the same inclusion and exclusion criteria to both datasets. We manually went through all collected posts from Rednote and Weibo and removed 340 of them, which were irrelevant to LLM fortunetelling. The removal criteria included advertisements, promotional content, and posts unrelated to fortunetelling practices with LLMs. Lastly, we had 1,045 posts, with 37 of them containing videos and 1,008 posts having images. The final dataset includes 859 posts from Rednote and 186 posts from Weibo, and all subsequent analyses report both aggregated results and platform-specific distributions. These posts averaged 620 likes and 66 comments per post. All data were publicly available at the time of collection. Following ethical research practices, we excluded private or non-public information, anonymized user identifiers, and focused exclusively on aggregated analyses to protect user privacy and minimize potential risks.

Coding the Posts

We adopted an open coding approach to understand LLM-based fortunetelling shared in the posts, following established procedures in related research (Zhang et al. 2023). A post could be assigned one code in each category and multiple codes within a category (*e.g.*, prompt structure). If a post did not belong to a code in a category, we mark it as None in that category. The collected posts contain videos or images, which makes it difficult to develop classifiers for automatic analyses. Therefore, we manually coded all

²<https://www.xiaohongshu.com/>

³<https://passport.weibo.com/>

⁴https://github.com/cv-cat/Spider_XHS

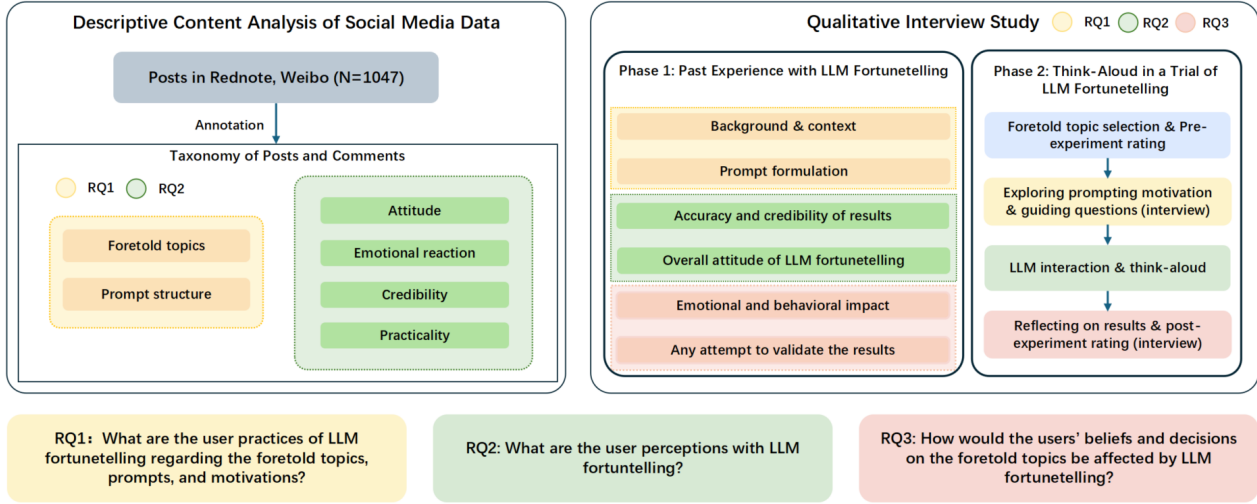


Figure 2: Research flow of understanding fortunetelling with LLMs by social media data analysis and interview study.

1,045 posts. Specifically, two authors annotated 266 posts to develop a codebook. They first independently reviewed the post content and assigned codes related to our research questions, resulting in two independent codebooks. The two authors then collaboratively produced a final codebook through three rounds of discussions and iterations. The two coders then re-coded the 266 posts based on the final codebook. Throughout the coding process, we consulted relevant literature (Zhang et al. 2023) and psychology research (Huang et al. 2024) to guide the definition of each code. For example, “Credibility” corresponds to the commonly used practicality evaluation dimensions in AI-assisted decision-making research (Ma et al. 2024a). We also identified a small number of posts discussing how LLM fortunetelling affected users’ decisions (RQ3). Therefore, the final codebook focused on the themes of user practices (RQ1, Table 1) and perceptions (RQ2, Table 2), with an additional theme about the post topics. To assess the reliability of the open coding process, we calculated inter-coder agreement for each code using Cohen’s Kappa, resulting in an overall mean of 0.845 (SD = 0.095). Specifically, within the “Foretold Topics” category (Table 1), the codes of “Lost Property” and “Fertility” have the highest agreement of $\kappa = 1$, while within the “Emotional Reaction” category (Table 2), the “Humorous Commentary” code has the lowest agreement of $\kappa = 0.712$. After establishing the taxonomy, the two coders worked together to annotate the remaining posts. Finally, we obtained a labeled dataset for analysis. In the following analysis, we primarily analyzed the merged dataset from the two platforms to capture overall patterns. As a preview of these comparisons (see Table 1 and Table 2), the distributions on Rednote and Weibo are generally comparable: although some differences exist (e.g., the “Career” topic: 24.10% vs. 16.67%), key categories appear on both platforms (e.g., the “General Fortune” category: 16.18% vs. 15.59%), and the overall attitude trends are similar (e.g., the “Positive” category: 58.32% vs. 72.04%). This supports

our use of the merged dataset for analysis while preserving platform-specific nuances.

Theme: Topics of Posts

Posts share three codes (Figure 1) about their topics, *i.e.*, discussion on fortunetelling (proportion for posts: 18.66%), sharing experiences with fortunetelling (74.83%), and sharing prompts (23.35%). Within the category of discussions on fortunetelling, users’ conversations further focused on specific aspects, including reliability and safety concerns (44.78% of discussion posts, *e.g.*, worries that AI simply tells people what they want to hear), model performance comparisons (30.43%, *e.g.*, debating whether Gemini or GPT provides more accurate predictions), suggestions for using LLM fortunetelling (8.70%, *e.g.*, encouraging others to use AI to avoid bad outcomes), suggestions for prompting LLMs (10.87%, *e.g.*, proposing to generate a birth chart first before asking the model), and other topics (5.22%, *e.g.*, reflections on how tools like DeepSeek affect people’s sense of destiny). When sharing fortunetelling experiences, posters share screenshots of their fortunetelling with LLM and describe the results generated by LLM and their feelings in the post content. And when sharing prompts, posters share their complete prompts for fortunetelling using LLM in the post content.

Theme: User Practices (RQ1)

As shown in Table 1, we identify multiple codes that fall into the following two categories regarding user practices with LLM fortunetelling.

Foretold topics. This category refers to the topics that users would like to foretell using LLMs. The most common topics include romance (27.08% of posts), career (22.78%), and general fortune (16.08%). Romance-related queries often concern relationships and future partners, while career-related queries focus on job changes and career development. General fortune-telling typically involves users pro-

viding personal information such as birth date and birthplace for holistic predictions about life, marriage, career, and wealth. The topic of academic studies (10.91%) is also notable. Users often consult LLMs about exam results or higher education opportunities. Wealth-related queries (9.86%) mainly focus not on financial planning, but on luck and prosperity. For example, users ask when they might become wealthy or whether their “fortune in wealth” is favorable. Some users even use LLM fortunetelling to locate lost objects (3.54%), showing its role in both major life concerns and everyday practical issues. These findings suggest that LLM fortunetelling combines traditional concerns such as romance and wealth with modern uncertainties related to education and career, functioning as a general resource for navigating everyday uncertainty.

Prompt structure. This category refers to the prompt structure used by users of LLM fortunetelling, including common character settings, output format, and so on (Google 2025). As shown in Table 5 in the Appendix, the prompt structure shared by users on social media is relatively fixed, combining strategies of role-playing LLMs with information needed in traditional fortunetelling practices. Most prompts involve contextual input, such as providing birth information and astrological information, often corresponding to specific traditional systems such as Bazi and Ziwei Doushu, as well as Western astrology. Our dataset shows that Chinese traditional methods dominate these practices, with Bazi/Four Pillars appearing in 215 posts (20.57% of all posts; 65.95% among posts that explicitly mention a method), followed by Ziwei Doushu (3.45%) and other systems such as Liuyao (3.25%). In contrast, Western methods such as astrology (5.55%) and tarot (2.20%) appear less frequently. Meanwhile, most prompts provide complete task instructions, such as “*Analyze my horoscope.*” Users also frequently use character settings, instructing the LLM to act as a professional fortune-teller, and knowledge recall (15.22%), citing classic texts or traditional methods. These practices place interactions with LLM within a familiar cultural framework, indicating that users view LLM as a continuation of traditional fortunetelling systems.

Theme: User Perceptions (RQ2)

Table 2 summarizes the user perceptions of LLM fortunetelling shared in the posts, which fall into the following categories.

General attitude. This category refers to the user’s overall attitude towards LLM fortunetelling, including positive, neutral, and negative attitudes. Most users (60.77% of posts) hold a positive attitude toward LLM fortunetelling in social media, often praising the accuracy of predictions or joking about LLM replacing human fortune-tellers. Neutral opinions are also common (16.84%), while explicitly negative views remain in the minority (6.32%), mainly criticizing inaccurate results or the LLM’s tendency to cater to users’ expectations. Overall, users generally engage with LLM fortunetelling positively while remaining aware of its limitations.

Emotional reactions. This category refers to the user’s emotional reaction (*e.g.*, surprise, worry) to the LLM for-

tunetelling results as reflected in the post. The most emotional reaction of users (34.45% of posts) to fortunetelling results is surprising when the prediction matches reality. Many users regard this coincidence as validation, thereby increasing their investment in the results. Additionally, some users interpret predictions humorously (14.45%). For example, they joke about the irrationality of the results, using a relaxed tone to deflect their seriousness. Other users (10.14%) use predictions as a source of psychological support, offering comfort or encouragement and hoping for a positive outcome. Meanwhile, a small number of users perceive contradictions, such as concerns about unfavorable marriage or career prospects (3.64%), demonstrating that LLM fortunetelling can sometimes trigger negative emotions. Overall, these emotional reactions reflect both users’ expectations and entertaining interpretations of the results, while also revealing their psychological impact.

Perceived credibility. This category captures how users assess the credibility of LLM fortunetelling results, ranging from trust to skepticism and outright distrust. Among posts, 36.08% convey rational skepticism, often asking if the system is saying nice things on purpose or deliberately catering to user expectations. Meanwhile, 28.42% of posts express a strong trust, with users reporting that predictions aligned with real-life experiences, such as receiving a job offer after being told their career luck is improving. A smaller proportion (6.22%) voices a clear distrust, often pointing out obvious inaccuracies, for example, complaining that the system could not even get the Bazi right. Overall, these findings suggest that while many users view LLM predictions as somewhat credible, a notable portion remains doubtful or critical.

Perceived practicality. Perceived practicality concerns whether users regard LLM fortunetelling as useful in everyday life. About half of the posts (50.24%) describe it as practical, noting applications such as finding lost items or guiding routine decisions, while 6.89% of posts criticize its lack of accuracy and question its real-world value. At the same time, 42.39% of posts do not explicitly address usefulness, reflecting that systematic reflections on utility are less frequent than personal anecdotes. Overall, users tend to perceive LLM fortunetelling as practical when it provides help in specific situations, but they rarely engage in broader or systematic evaluations of its overall value.

Qualitative Interview Study with Users of LLM Fortunetelling

While the descriptive content analysis of social media data provided a comprehensive taxonomy of users’ practices (RQ1) and perceptions (RQ2) of LLM fortunetelling, we still lack in-depth insights into their underlying reasons and the impact on users’ decision-making behaviors (RQ3). Therefore, we further conducted semi-structured interviews with 20 users of LLM-based fortunetelling, focusing on how they interpret and respond to the fortunetelling results. (Liu et al. 2025)

Participants

Table 1: Distribution of foretold topics in LLM fortune-telling, identified through a content analysis of 1,045 social media posts (859 from Xiaohongshu and 186 from Weibo), with the specific quantities and frequencies for each platform.

Categories	Codes	Sample	Total	Xiaohongshu	Weibo
Foretold Topics	Career	“The readings I did for myself and my friends about our current work situations were quite accurate, and even the predictions about career development were somewhat relevant.”	238 (22.78%)	207 (24.10%)	31 (16.67%)
	Romance	“It says I am skeptical about marriage and will marry late or not marry at all.”	283 (27.08%)	256 (29.80%)	27 (14.52%)
	General Fortune	“It even got my Bazi wrong.”	168 (16.08%)	139 (16.18%)	29 (15.59%)
	Health	“It’s actually spot on, I can’t believe it even knew my right leg keeps getting injured!”	48 (4.59%)	44 (5.12%)	4 (2.15%)
	Academic Studies	“Chatgpt calculated that I would be admitted and would likely receive a scholarship. Turns out I was the only one in my major who received a scholarship.”	114 (10.91%)	108 (12.57%)	6 (3.23%)
	Fertility	“It was right when he said I’d have a daughter.”	25 (2.39%)	25 (2.91%)	0 (0.00%)
	Wealth	“It said I’ll make 150 million in middle age, I don’t even dare to believe it!”	103 (9.86%)	86 (10.01%)	17 (9.14%)
	Interpersonal Relationships	“The result from the AI says I need to watch out for bad people this year.”	38 (3.64%)	33 (3.84%)	5 (2.69%)
	Lost Property	“I finally found the ring! DeepSeek, I’ll stand by you, it said near the water source and I ended up finding it in the kitchen sink.”	37 (3.54%)	36 (4.19%)	1 (0.54%)
	Other Topic	“I used DeepSeek for fortune-telling to choose which blind box to buy, and it actually picked the one I wanted!”	91 (8.71%)	85 (9.90%)	6 (3.23%)
	None	-	204 (19.52%)	121 (14.09%)	83 (44.62%)

We recruited 7 creators of our collected Rednote posts who accepted our invitation sent via private messages (we sent 67 invitations in total). We also used two principles to select 13 out of 100 respondents to our online advertisements in the authors’ social networks (*e.g.*, via WeChat group chats). First, they should have experience in LLM fortunetelling. Second, their occupations or majors and foretold topics in the past should be diverse. In the invitations or advertisements, individuals would receive a detailed overview of the study, including its purpose, what participation involves, the expected duration of their involvement, any potential risks and benefits, and their rights as participants. After confirming the consent, participants would engage in a one-hour semi-structured online interview. Each participant received a cash compensation of 50 Chinese Yuan (equivalent to 6.92 USD) after finishing the interview. Table 4 in the Appendix describes the detailed information of our 20 participants (P1-20, 15 females and 5 males, age range: 20-43).

Procedure

The interview study consisted of two phases: one asks about the participants’ past experiences with LLM fortunetelling; the other seeks to learn their in-situ thoughts in a trial of LLM fortunetelling. Table 3 in the Appendix summarizes the questions we asked in two phases, respectively.

Phase 1: Past Experience with LLM Fortunetelling To get started, participants are asked to recall their previous experiences using LLMs for fortunetelling. We ask their foretold topics, motivations, prompting strategies (RQ1), perceptions of the results (RQ2), and their subsequent thoughts

and behaviors (RQ3). Key questions include the context in which participants used LLM fortunetelling, their prompt formulation, immediate reactions to the LLM responses, perceived accuracy and credibility, and overall attitudes. Specifically, we have questions asking about the emotional and behavioral impact of the fortunetelling results on participants, as well as any attempt to validate the results.

Phase 2: Think-Aloud in a Trial of LLM Fortunetelling

Following phase 1, participants select a foretold topic (*e.g.*, romance, academic studies) that is closely related to them at that moment to conduct LLM (mostly DeepSeek) fortunetelling with LLMs on their computers. We prepare example prompts sourced from the related posts on social media platforms for reference. This trial can provide a fresh experience of LLM fortunetelling to understand how participants interpret and are affected by the results. Specifically, before and after the trial, we ask participants to rate their likelihood of making a specific decision (*e.g.*, quit a job or not) or level of confidence towards the foretold topics in the future (*e.g.*, paper acceptance) from 0 to 100%. Prior to the trial, we ask participants about their prompting motivation and present them with guiding questions asking about their immediate emotional reactions to the LLM answers (*e.g.*, feeling understood, encouraged, or confused), perceived strengths or shortcomings of the answers, and potential behavioral impact. During the trial, participants are asked to think aloud, expressing their thoughts related to the questions above in real time. All interactions are screen- and audio-recorded, anonymized, and treated confidentially. For the interview recordings, two of the authors transcribed them

Table 2: This table systematically categorizes user perceptions of LLM-generated fortunetelling results into five key dimensions: Attitude, Emotional Reaction, and Assessments of Credibility and Practicality. It provides both qualitative examples and descriptive statistics (from 1,045 posts, 859 Xiaohongshu, 186 Weibo), specifically including platform-specific quantities and frequencies to offer a comprehensive overview of the user perceptions of the fortunetelling results given by LLMs (RQ2).

Categories	Codes	Sample	Total	Xiaohongshu	Weibo
Attitude	Positive	“The AI’s too good at this, feels like fortune tellers are about to lose their jobs.”	635 (60.77%)	501 (58.32%)	134 (72.04%)
	Neutral	“Some of the AI’s fortunetelling is spot on, but some parts are way off.”	176 (16.84%)	130 (15.13%)	46 (24.73%)
	Negative	“Not sure who’s hyping this up, I checked it out and the fortunetelling was way off.”	66 (6.32%)	60 (6.98%)	6 (3.23%)
	None	“DeepSeek says I’ll meet my true love this year.”	167 (15.98%)	167 (19.44%)	0 (0.00%)
Emotional Reaction	Surprise	“I asked DeepSeek what place I’d get in the interview, and it said second — and guess what, I really came in second!”	360 (34.45%)	310 (36.09%)	50 (26.88%)
	Worry	“AI predicted that my husband and I are destined to get divorced. What should I do?”	38 (3.64%)	33 (3.84%)	5 (2.69%)
	Humorous Commentary	“The prediction says I’ll meet my true love this year—great! Can you also calculate his phone number for me?”	151 (14.45%)	128 (14.90%)	23 (12.37%)
	Hope and Encouragement	“The prediction says that he will definitely reach out to me between 1 and 2 p.m. this afternoon—I hope it comes true!”	106 (10.14%)	65 (7.57%)	41 (22.04%)
	Contradiction	“The AI says I’ll have four kids and the first will be a girl. . . but my first is a boy, and I’m not having any more!”	53 (5.07%)	45 (5.24%)	8 (4.30%)
	Other Reactions	“I’ll just wait and see whether this result turns out to be accurate or not.”	16 (1.53%)	16 (1.86%)	0 (0.00%)
	None	-	322 (30.82%)	262 (30.50%)	60 (32.26%)
Credibility	Skeptical	“Why do I feel like it’s saying nice things on purpose?”	377 (36.08%)	313 (36.44%)	64 (34.41%)
	Trust	“It’s really accurate, he said I’d get an offer this month and I started my job today.”	297 (28.42%)	264 (30.73%)	33 (17.74%)
	Distrust	“It can’t even get the Bazi right.”	65 (6.22%)	58 (6.75%)	7 (3.76%)
	None	-	303 (29.00%)	222 (25.84%)	81 (43.55%)
Practicality	Practical	“Last time I lost something, I told it the rules for the calculation and it helped me find it, that’s really awesome.”	525 (50.24%)	462 (53.78%)	63 (33.87%)
	Impractical	“Not accurate, it said I could score 638 but I only got 566.”	72 (6.89%)	66 (7.68%)	6 (3.23%)
	None	-	443 (42.39%)	328 (38.18%)	115 (61.83%)

into text and conducted a thematic analysis. They first familiarized themselves by reviewing all the text scripts independently. After several rounds of coding with comparison and discussion, they finalized the codes of all the interview data. We counted the occurrences of codes and incorporated these qualitative findings in the following presentation of our results.

Findings

Motivations (RQ1): decision support, emotional comfort, entertainment, and social substitution. Users turned to LLM fortune-tellers when making critical decisions, such as furthering their studies (P7) or resignation (P9), seeking guidance and reassurance. Meanwhile, participants frequently emphasized the dimension of emotional comfort by using LLM fortunetelling, *e.g.*, to seek support before exams or competitions (P1, P18) or after conflicts in close relationships (P4, P6). In fact, many participants mentioned that they primarily tried fortunetelling out of curiosity or entertainment. Some users also used LLM fortunetelling to confirm pre-existing thoughts (P17) or as a substitute for friends

or counselors (P16), fulfilling needs for communication and companionship. One participant (P11) even described consulting LLM fortunetelling as a habitual step before making decisions.

Participants exhibited highly different perceptions of LLM fortunetelling results, which are often presented in a partially truthful narrative style (RQ2). In the experiment, some users found the results reliable, particularly convincing when using horoscopes to accurately predict the timing of reconciliation with a boyfriend (P10) or academic milestones (P13). A few results closely matched their personal experiences, *e.g.*, the family experiencing a flood (P5) and exchanging programs in Hong Kong (P8), leaving users shocked. However, many users (P2, P3, P7, P10, P14, P17, P18) found the results vague and general without specific operational guidance, making them still feel uncertain about their decisions. During the trials of LLM fortunetelling during the study, users often found themselves with a sense of partial conviction and partial skepticism. They were surprised by the corrected information in LLM’s responses, but they also identified obvious errors or ambiguous language.

“When I noticed that the AI miscalculated my age, I felt it was foolish.” (P18)

Participants have neutral or positive attitudes towards LLM fortunetelling but are aware of the LLM’s tendency to cater to their expectations (RQ2). Some participants pointed out potential concerns of LLM fortunetelling, such as over-reliance (e.g., relying too heavily on LLM fortunetelling results for important decisions) and privacy issues (e.g., disclosing birth data or sensitive personal experiences). Nevertheless, most users regarded LLM fortunetelling as harmless and primarily treated it as a form of entertainment or psychological comfort. In the trials during the study, seventeen users found that the LLM often tailors its responses to align with their tone and expectations. For example, the responses often contained compliments (e.g., *“You’ll achieve great success in your future career”*) or initially presented challenges but ultimately concluded positively (e.g., *“You may face some setbacks right now, but things will soon turn around”*). Interestingly, P9 and P19 said that they expected or readily accepted LLM’s catering in fortunetelling, as it fostered confidence and hope, alleviating anxiety and uncertainty in the short term. Conversely, P4 and P7 preferred the LLM to maintain a neutral, rational, or objective tone to give fortunetelling results, saying that this can enhance the results’ credibility and reference value.

LLM fortunetelling exerts modest influences on users’ beliefs and decision-making processes (RQ3). Our findings suggest that LLM fortunetelling mainly provides psychological reassurance and modestly influences the timing of decisions rather than overall directions (Ma et al. 2024b).

Influence on beliefs and emotions. Several participants (P1, P10, P11, P12, P19) described using LLM fortunetelling as a form of psychological adjustment. It offers reassurance in times of stress or anxiety, and provides narratives that resonate with their inner struggles. For instance, P11 consulted the LLM when feeling miserable at work, and upon receiving a reframe—that the feeling stemmed from a longing for freedom constrained by reality, and that work could be approached as an “observation game”—found the perspective genuinely helpful and reported feeling less drained by office politics thereafter. Additionally, in the experiment, 17 participants chose a topic to predict and rated their confidence in predictive outcomes generated by LLMs, covering both broad life domains (e.g., romance, career) and specific personal goals (e.g., whether she could successfully lose weight). As shown in the Table 4 in the Appendix, five participants showed no change before and after the interaction, while the remaining exhibited an average change of 9.53 points (SD = 8.84), with the largest shift being 30 points and the smallest 0. These results suggest that LLM fortunetelling can calibrate the degree of certainty in participants’ beliefs about the future.

Influence on decisions. By contrast, the direct influence of LLM fortunetelling on concrete decision-making is limited. While most participants emphasized that they relied primarily on personal judgment, financial conditions, and family considerations, the advice generated by LLM some-

times shaped the timing of actions. For example, participants reported delaying resignation (P9), postponing competitions (P18), or suspending investments (P16) based on LLM suggestions. P9, who was already strongly considering leaving their job, consulted the LLM about career prospects and upon being advised not to quit in 2025—an output that matched existing worries—chose to hold on. Rather than changing decision directions, LLM fortunetelling mainly reinforced participants’ existing inclinations and influenced the timing of actions. Participants often felt reassured when the generated results aligned with their preferred choices. For example, P17 entered the consultation already leaning toward staying in the current department, and the fortunetelling result unexpectedly aligned with this inclination, reinforcing the sense that remaining might be the more reasonable choice. This shows that LLM fortunetelling did not necessarily change participants’ decision directions but strengthened their beliefs in the legitimacy of their own preferences. Importantly, these effects do not indicate changes in decision outcomes, but rather influence the timing and pacing of actions.

Users actively interpret, filter, and internalize LLM-generated results, and these practices influence how fortunetelling shapes their beliefs about everyday life (RQ3). Rather than passively accepting results, participants reinterpret vague predictions as actionable or psychologically useful guidance. For example, P12 transformed “pay attention to your kidneys” into “stay up less late”. At the same time, P11 summarized general health warnings as reminders to monitor their diet. Some even integrated generic outputs into specific plans, such as aligning career (P9) or exam schedules (P19) with the predicted time. The impact is therefore selective. Short-term and concrete predictions can foster small belief adjustments and motivate behavioral fine-tuning. In contrast, long-term or ambiguous predictions are approached with skepticism and mainly provide psychological reassurance.

Discussion

In this work, we offer comprehensive findings on user practices, perceptions, and impacts of LLM fortunetelling via analyzing social media data and interviewing users. This section discusses the implications of these findings for related research fields, values and risks of LLM fortunetelling, as well as limitations and future work.

Theoretical Explanation of Findings

Our findings on LLM fortunetelling can be theoretically explained through the lenses of motivated reasoning, confirmation bias, and the Barnum effect. Users approached LLM fortunetelling driven by desires for reassurance, validation, and emotional comfort. This motivational basis reflects motivated reasoning (Kunda 1990), which posits that individuals tend to interpret information in a way that aligns with their own desires and emotional needs—this explains why users consistently seek outcomes that conform to rather than contradict their inner inclinations in various matters, from career transitions to interpersonal relationship choices.

The ways in which users perceived and interpreted LLM outputs further reveal confirmation bias (Nickerson 1998) and the Barnum effect (Huizi, Yiwei et al. 2022) at work. Participants selectively associated vague predictions with personal experiences and reinterpreted ambiguous outputs as actionable guidance. This resonates with the Barnum effect, where generic statements feel personally meaningful precisely because individuals project their own experiences onto them. At the same time, confirmation bias led users to favor outputs that aligned with their prior beliefs—a tendency that manifests across seeking, interpreting, and recalling information, and is especially pronounced among individuals with stronger prior beliefs (Boonprakong et al. 2025). Notably, seventeen participants observed that the LLM tailored its responses to match their tone and expectations, yet largely accepted this catering: P9 and P19 even anticipated and welcomed it, as it fostered confidence and alleviated anxiety. This pattern echoes findings in human-AI collaboration, where AI outputs that mirror users' initial judgments boost users' confidence in those judgments and amplify confirmation bias (Rosbach et al. 2025).

Because users selectively interpreted LLM outputs to align with their existing beliefs, the influence on their actual decisions was modest but consistent. Rather than altering decision directions, LLM fortunetelling primarily reinforced existing inclinations—P9 held off on resignation because the LLM output matched existing worries. Users adjusted the timing of actions when outputs aligned with prior preferences, and overall confidence in existing beliefs shifted by an average of 9.53 points. This pattern resembles positive illusions (Taylor and Brown 1988), where selectively internalized information sustains a sense of control over uncertain futures. Taken together, LLM fortunetelling functions less as a predictive tool and more as a space for emotional reflection and meaning-making (Zheng et al. 2025), reconfiguring traditional fortunetelling into a hybrid cultural practice that combines belief reinforcement with LLM-mediated emotional support for navigating everyday uncertainty.

Platform Dynamics and Social Implications of LLM Fortunetelling

LLM fortunetelling operates not only as an individual interaction with LLM, but also as a platform-mediated social practice shaped by social media. By sharing LLM-generated fortunes in public posts, users turn private uncertainty into visible narratives that circulate within platform environments, allowing fortunetelling experiences to be encountered and interpreted by broader audiences. Social media platforms amplify certain fortunetelling narratives, particularly those that are emotionally resonant or perceived as coincidentally “accurate.” Features such as algorithmic feeds (Jose et al. 2025; Mousavi, Gummedi, and Zannettou 2024) increase the visibility of these cases, which may contribute to a skewed impression of the reliability of LLM-generated fortunes, even when individual users remain skeptical. Such visibility may contribute to the perception that LLM fortunetelling is a socially acceptable way of seeking reassurance under uncertainty.

The public framing of fortunetelling may also blur the boundary between entertainment and meaningful guidance. The coexistence of joking and serious use within the same social media context allows users to oscillate between detachment and belief, while experiencing the accumulated emotional and cognitive impact. From a platform perspective, LLM fortunetelling is associated with how uncertainty is discussed and shared in online contexts, pointing to potential implications beyond individual interactions.

Ambiguity, Ingratiation, and Risk in LLM-Based Fortunetelling

Our findings suggest that users often underestimate the risks associated with LLM-based fortunetelling. While some participants expressed concerns about dependence or privacy breaches, many overlooked the longer-term consequences of misleading outputs, repeated reliance, and the disclosure of sensitive personal information. This echoes prior work showing that users frequently struggle to recognize algorithmic harms in everyday interactions (Wu et al. 2025). One factor that may contribute to these risks is the ambiguous and ingratiating nature of LLM-generated fortunetelling outputs. Consistent with previous research (Han et al. 2025), these responses tend to rely on vague but affirmative language to create a feeling of being understood or accurately “predicted,” even without factual grounding. While such responses may provide emotional reassurance during uncertainty, they may also reduce critical scrutiny and encourage overestimation of reliability, particularly when users repeatedly reference divination results.

Current mitigation strategies, such as brief disclaimers (e.g., “for reference only”), may have limited effectiveness. Prior work (Bo, Wan, and Anderson 2025) suggests that mitigating over-reliance requires supporting users' reflective engagement with LLM-generated guidance. Specifically, these interventions might include embedded reflective prompts (e.g., “What other factors are important in your decision-making?”) and structured journaling features to invite users to record their thoughts along with the prediction. Additionally, contextual framing cues (e.g., differentiating between “entertainment” and “decision support” modes) and adaptive feedback (e.g., showing how the prediction affects user confidence or emotions) could further encourage critical reflection. We also suggest clearer warnings and privacy prompts when users share sensitive information. Beyond interface design, technological solutions can be explored, such as fine-tuning models or integrating prompt-based detection mechanisms to automatically identify sensitive or privacy-related user inputs, and applying filtering or rewriting modules to sanitize such inputs before inference (Biswas and Talukdar 2026).

Taken together, these findings highlight a paradox in LLM fortunetelling: while emotionally supportive outputs may reassure users, repeated reliance on such advice may also weaken careful scrutiny.

Limitations and Future Work

This study has several limitations. First, our social media data is limited to public posts on RedNote and Weibo, which

may miss private or sensitive use scenarios. As a result, users who engage in LLM fortunetelling without publicly sharing their experiences are not captured, potentially introducing sampling bias toward more expressive or socially motivated users. Second, although we pool data from RedNote and Weibo to identify overall patterns of LLM fortunetelling, platform-specific affordances may still shape user behavior. In addition, our interview sample of 20 participants may not fully represent demographic diversity and may be subject to self-selection bias, as participation was voluntary and participants were partly recruited through the authors' networks. Future research could expand the sample size, recruit more diverse participants, and explore semi-private interactions to capture a broader range of experiences. Additionally, our social media data captures only observable behaviors, and our interviews provide a single-point, self-reported perspective (e.g., short-term confidence ratings). As a result, we cannot assess the long-term effects of LLM fortunetelling on users.

Moreover, our analysis focuses on posts as acts of public self-expression rather than interactional dynamics in comment threads. In addition, our approach relies on descriptive content analysis and does not incorporate computational modeling. Future work can further conduct quantitative analyses (e.g., regression) to examine the relationship between the content of LLM fortunetelling posts and the types of comments they receive. Such analyses could deepen our understanding of how different content features shape user engagement and responses. Finally, this study examines Chinese users and fortunetelling mostly using Chinese traditional terminology. Cultural factors may affect users' perceptions and practices of LLM fortunetelling. Future research should explore other cultural contexts to assess the broader applicability of the results.

Conclusion

We investigate how people engage in fortunetelling with LLMs through two qualitative studies, combining a descriptive content analysis of 1,045 posts on Chinese social media with interviews of 20 LLM fortunetelling users. Our findings suggest that users turn to LLM fortunetelling for emotional support and decision-related reflection, and that such interactions may shape users' beliefs and decision-making in subtle ways. Overall, this work positions fortunetelling as a meaningful form of daily interaction with LLMs. Finally, this study highlights the potential impact of LLM fortunetelling, particularly related to users' emotional experience, interpretation of the results, and perceptions of associated risks.

References

- Biswas, A.; and Talukdar, W. 2026. Guardrails for trust, safety, and ethical development and deployment of Large Language Models (LLM). *arXiv preprint arXiv:2601.14298*.
- Bo, J. Y.; Wan, S.; and Anderson, A. 2025. To Rely or Not to Rely? Evaluating Interventions for Appropriate Reliance on Large Language Models. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.
- Boonprakong, N.; Pareek, S.; Tag, B.; Goncalves, J.; and Dingler, T. 2025. Assessing Susceptibility Factors of Confirmation Bias in News Feed Reading. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.
- Buçinca, Z.; Swaroop, S.; Paluch, A. E.; Doshi-Velez, F.; and Gajos, K. Z. 2025. Contrastive explanations that anticipate human misconceptions can improve human decision-making skills. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1–25.
- Chen, S.; Xie, J.; Wang, G.; Wang, H.; Cheng, H.; and Huang, Y. 2025. From Scores to Careers: Understanding AI's Role in Supporting Collaborative Family Decision-Making in Chinese College Applications. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.
- Cheng, S.; Han, M.; Lu, Z.; Xu, W.; and Yu, J. 2025. Beat ChatGPT at Fortune-Telling—An Attempt to Optimize a Large Language Model. *Applied and Computational Engineering*, 100: 194–208.
- Cho, H.; Seo, J. A.; Lee, J.; Kim, C.-M.; and Nam, T.-J. 2025. ShamAln: Designing Superior Conversational AI Inspired by Shamanism. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.
- Danry, V.; Pataranutaporn, P.; Groh, M.; and Epstein, Z. 2025. Deceptive Explanations by Large Language Models Lead People to Change their Beliefs About Misinformation More Often than Honest Explanations. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.
- Google. 2025. Prompt Engineering for Generative AI. <https://developers.google.com/machine-learning/resources/prompt-eng>. Last updated 2025-08-25.
- Grall-Bronnec, M.; Bulteau, S.; Victorri-Vigneau, C.; Bouju, G.; and Sauvaget, A. 2015. Fortune telling addiction: Unfortunately a serious topic about a case report. *Journal of Behavioral Addictions*, 4(1): 27–31.
- Han, P.; Kocielnik, R. D.; Song, P.; Debnath, R.; Mobbs, D.; Anandkumar, A.; and Alvarez, R. M. 2025. Tracing Human-like Traits in LLMs: Origins, Real-World Manifestation, and Controllability. In *2nd Workshop on Models of Human Feedback for AI Alignment*.
- Hou, H.; Leach, K.; and Huang, Y. 2024. ChatGPT Giving Relationship Advice—How Reliable Is It? In *Proceedings of the International AAAI Conference on Web and Social Media*, volume 18, 610–623.

- Huang, J.-t.; Lam, M. H.; Li, E. J.; Ren, S.; Wang, W.; Jiao, W.; Tu, Z.; and Lyu, M. R. 2024. Apathetic or empathetic? evaluating llms' emotional alignments with humans. *Advances in Neural Information Processing Systems*, 37: 97053–97087.
- Huizi, C.; Yiwei, X.; et al. 2022. Understanding the Impact of Barnum Effect in Astrology: An Eye-tracking Study. *The Frontiers of Society, Science and Technology*, 4(12).
- Jose, J.; Geeng, C.; Morales, K. O.; McCoy, D.; and Greenstadt, R. 2025. What's in a Label? Propaganda Labels and User Sharing Behavior on Social Media Platforms. *Proceedings of the International AAAI Conference on Web and Social Media*, 19(1): 918–934.
- Kim, K.; Song, H.; Ryu, J.; Oh, C.; and Suh, B. 2025. BleacherBot: AI Agent as a Sports Co-Viewing Partner. 10.1145/3706598.3713259. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.
- Kim, S.; Razi, A.; Alsoubai, A.; Wisniewski, P. J.; and De Choudhury, M. 2024. Assessing the impact of online harassment on youth mental health in private networked spaces. In *Proceedings of the international AAAI conference on web and social media*, volume 18, 826–838.
- Kunda, Z. 1990. The case for motivated reasoning. *Psychological bulletin*, 108(3): 480.
- Li, Z.; Zhu, H.; Lu, Z.; Xiao, Z.; and Yin, M. 2025. From Text to Trust: Empowering AI-assisted Decision Making with Adaptive LLM-powered Analysis. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1–18.
- Liu, M.; Wang, T.; Cohen, C. A.; Li, S.; and Xiong, C. 2025. Understand User Opinions of Large Language Models via LLM-Powered In-the-Moment User Experience Interviews. *arXiv preprint arXiv:2502.15226*.
- Liu, P.; Yuan, W.; Fu, J.; Jiang, Z.; Hayashi, H.; and Neubig, G. 2023. Pre-train, prompt, and predict: A systematic survey of prompting methods in natural language processing. *ACM computing surveys*, 55(9): 1–35.
- Ma, S.; Wang, X.; Lei, Y.; Shi, C.; Yin, M.; and Ma, X. 2024a. “Are You Really Sure?” Understanding the Effects of Human Self-Confidence Calibration in AI-Assisted Decision Making. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, CHI '24. New York, NY, USA: Association for Computing Machinery. ISBN 9798400703300.
- Ma, S.; Wang, X.; Lei, Y.; Shi, C.; Yin, M.; and Ma, X. 2024b. “Are you really sure?” Understanding the effects of human self-confidence calibration in AI-assisted decision making. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 1–20.
- Meier, V. 2022. Fortunetelling As a Fraudulent Profession? The Gendered Antigypsyist Motif of Fortunetelling and Persecution by the Criminal Police. *Critical Romani Studies*, 5(1): 30–48.
- Mousavi, S.; Gummadi, K. P.; and Zannettou, S. 2024. Auditing algorithmic explanations of social media feeds: A case study of tiktok video explanations. In *Proceedings of the International AAAI Conference on Web and Social Media*, volume 18, 1110–1122.
- Nickerson, R. S. 1998. Confirmation bias: A ubiquitous phenomenon in many guises. *Review of general psychology*, 2(2): 175–220.
- Rosbach, E.; Ammeling, J.; Krügel, S.; Kießig, A.; Fritz, A.; Ganz, J.; Puget, C.; Donovan, T.; Klang, A.; Köller, M. C.; et al. 2025. “When Two Wrongs Don't Make a Right”—Examining Confirmation Bias and the Role of Time Pressure During Human-AI Collaboration in Computational Pathology. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1–18.
- Salimzadeh, S.; He, G.; and Gadiraju, U. 2024. Dealing with uncertainty: Understanding the impact of prognostic versus diagnostic tasks on trust and reliance in human-AI decision making. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, 1–17.
- Skryabin, V. Y. 2020. An addiction to seeking fortune-telling services: a case report. *Journal of Addictive Diseases*, 38(2): 223–228.
- Smith, R. J. 2019. An Overview of Chinese Fortune Telling in Traditional Times. *The Routledge Encyclopedia of Traditional Chinese Culture*, 297–320.
- Tan, X.; Van Prooijen, J.-W.; and van Lange, P. A. 2022. Positive fortune telling enhances men's financial risk taking. *Plos one*, 17(9): e0273233.
- Taylor, S. E.; and Brown, J. D. 1988. Illusion and well-being: a social psychological perspective on mental health. *Psychological bulletin*, 103(2): 193.
- Tlili, A.; Shehata, B.; Adarkwah, M. A.; Bozkurt, A.; Hickey, D. T.; Huang, R.; and Agyemang, B. 2023. What if the devil is my guardian angel: ChatGPT as a case study of using chatbots in education. *Smart learning environments*, 10(1): 15.
- Tsai, C.-H.; You, Y.; Gui, X.; Kou, Y.; and Carroll, J. M. 2021. Exploring and promoting diagnostic transparency and explainability in online symptom checkers. In *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, 1–17.
- Wu, Y.; Dong, Y.; Mou, Y.; and Kim, K. J. 2025. How the Algorithmic Transparency of Search Engines Influences Health Anxiety: The Mediating Effects of Trust in Online Health Information Search. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.
- Zhang, X.; and Wang, Y. 2025. Swiping through the stars: exploring identity and intimacy through astrology sensemaking on Douyin. *new media & society*, 14614448251344588.
- Zhang, Y.; He, C.; Wang, H.; and Lu, Z. 2023. Understanding Communication Strategies and Viewer Engagement with Science Knowledge Videos on Bilibili. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, CHI '23. New York, NY, USA: Association for Computing Machinery. ISBN 9781450394215.

Zheng, X.; Li, Z.; Gui, X.; and Luo, Y. 2025. Customizing Emotional Support: How Do Individuals Construct and Interact With LLM-Powered Chatbots. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI '25. New York, NY, USA: Association for Computing Machinery. ISBN 9798400713941.

Ethics Checklist

1. For most authors...

- (a) Would answering this research question advance science without violating social contracts, such as violating privacy norms, perpetuating unfair profiling, exacerbating the socio-economic divide, or implying disrespect to societies or cultures? Yes. The social media data used in this study are anonymized and processed to protect user privacy, and no personally identifiable information is analyzed. Participants involved in interviews or surveys are treated with confidentiality and informed consent, ensuring their privacy is fully protected. The study avoids unfair profiling or cultural bias and does not contribute to social or economic inequality.
 - (b) Do your main claims in the abstract and introduction accurately reflect the paper's contributions and scope? Yes. For example, our main claims focus on the LLM fortunetelling, which accurately reflect the paper's scope.
 - (c) Do you clarify how the proposed methodological approach is appropriate for the claims made? Yes. We provide motivation, goal, or references for each used approach.
 - (d) Do you clarify what are possible artifacts in the data used, given population-specific distributions? Yes. The study acknowledges potential artifacts arising from population-specific distributions in the data. To mitigate this, the sample is designed to include participants from diverse age groups, occupations, and backgrounds. However, despite these efforts, some bias related to data availability and platform usage patterns may remain and cannot be completely avoided.
 - (e) Did you describe the limitations of your work? Yes.
 - (f) Did you discuss any potential negative societal impacts of your work? Yes.
 - (g) Did you discuss any potential misuse of your work? Yes. We discuss the potential misuse of LLM fortunetelling in the section "discussion".
 - (h) Did you describe steps taken to prevent or mitigate potential negative outcomes of the research, such as data and model documentation, data anonymization, responsible release, access control, and the reproducibility of findings? Yes. The study clearly describes multiple steps taken to prevent or mitigate potential negative outcomes. These include thorough data anonymization, careful data and model documentation, controlled access to sensitive resources, and responsible release of research outputs.
 - (i) Have you read the ethics review guidelines and ensured that your paper conforms to them? Yes.
2. Additionally, if you are using existing assets (e.g., code, data, models) or curating/releasing new assets, **without compromising anonymity...**
- (a) If your work uses existing assets, did you cite the creators? Yes. We cite them in reference or urls to the official websites.
 - (b) Did you mention the license of the assets? No. However, we have confirmed that all the assets in our paper can be used for academic purposes.
 - (c) Did you include any new assets in the supplemental material or as a URL? No. This study does not include any new assets in the supplemental material or as external URLs. All materials used are either described within the paper or based on existing, publicly available resources, ensuring clarity and completeness without introducing additional external assets.
 - (d) Did you discuss whether and how consent was obtained from people whose data you're using/curating? Yes. The study states that the data were collected from publicly available social media content in compliance with platform policies, and therefore individual consent was not explicitly obtained. To mitigate ethical concerns, all data were anonymized and analyzed only in aggregate for research purposes.
 - (e) Did you discuss whether the data you are using/curating contains personally identifiable information or offensive content? Yes. The paper clarifies that the data do not contain personally identifiable information, as all user identifiers are removed during preprocessing. Potentially offensive content is filtered or handled cautiously, and the analysis focuses on aggregated patterns rather than individual expressions.
 - (f) If you are curating or releasing new datasets, did you discuss how you intend to make your datasets FAIR? Yes. The paper discusses that the curated dataset follows FAIR principles by ensuring clear documentation and metadata for findability and reusability, standardized data formats for accessibility and interoperability, and well-defined usage conditions to support responsible reuse.
 - (g) If you are curating or releasing new datasets, did you create a Datasheet for the Dataset? Yes.
3. Additionally, if you used crowdsourcing or conducted research with human subjects, **without compromising anonymity...**
- (a) Did you include the full text of instructions given to participants and screenshots? Yes. The instructions provided to participants are summarized in the paper in a way that preserves anonymity.
 - (b) Did you describe any potential participant risks, with mentions of Institutional Review Board (IRB) approvals? The local institutions of the authors do not require nor provide ethical reviews for conducting human-subjects studies like us, which are not in med-

ical domains. Nevertheless, we took several ways to mitigate the ethical concerns.

- (c) Did you include the estimated hourly wage paid to participants and the total amount spent on participant compensation? Yes. Information about participant compensation, including estimated hourly wage and total cost, is reported.
- (d) Did you discuss how data is stored, shared, and deidentified? Yes. The study describes how data are securely stored, deidentified, and shared only under controlled conditions.

Acknowledgments of the Use of AI

We used AI (in particular large language models) for the following: supporting participants to experience LLM fortelling the during the qualitative interview study. Details can be found in the relevant sections. Authors take responsibility for the output and use of AI in this paper.

Appendix

Table 3: Interview questions in Phase 1 and Phase 2 and their corresponding research questions.

Question Content	Corresponding RQ
Phase 1: Past Experience with LLM Fortunetelling	
Have you ever asked LLM to provide you with important advice (e.g., on love, academics, or career) or use LLM for fortunetelling?	RQ1
Regarding your experience of LLM fortunetelling:	
What was the background (reason) at that time?	RQ1
How was LLM’s response?	RQ2
What was your intuitive feeling about the result given by LLM? Did you think it was highly credible?	RQ2
Do you think LLM fortunetelling had an impact on your subsequent mindset or thoughts? Can you give an example?	RQ3
Did these results make you rethink your choices, pause your original plans, or take new actions?	RQ3
Did you verify whether the AI fortunetelling was accurate afterward? What was the result?	RQ2
In what situations do you usually want to seek fortunetelling help? Is it for making decisions, emotional comfort, purely out of curiosity or entertainment?	RQ1
When you usually use LLM for fortunetelling, how do you organize your questions (prompts)? For example, do you refer to certain traditional fortunetelling terms?	RQ1
Where do you get these terms? For instance, from social media searches, or do you come up with them yourself?	RQ1
What is your overall attitude toward LLM fortunetelling—positive, neutral, or negative?	RQ2
Does it pose any harm to your life?	RQ2
What do you think is the source of LLM fortunetelling’s accuracy? Is it more like psychological suggestion, probability analysis, “actually a bit effective,” or something else?	RQ2
What words would you use to describe the experience brought by LLM fortunetelling? (e.g., reassuring, fun, understood, confused, distrustful, etc.)	RQ2
If LLM fortunetelling gives a conclusion you are unwilling to accept, how will you deal with it? Will you ignore it, adjust yourself, or seek a new interpretation?	RQ3
Phase 2: Think-Aloud in a Trial of LLM Fortunetelling	
Experiment Pre-/Post-Scoring (e.g., decision probability for “changing jobs”; confidence scores for love/career/health)	RQ3
What kind of answer do you hope to get from this fortunetelling? (e.g., emotional comfort, decision advice, curiosity satisfaction)	RQ1
What’s your first reaction to the answer? (e.g., shocked by accuracy, comforted, encouraged, advised, confused, etc.)	RQ2
What makes you feel “accurate” or “unreliable”? Why?	RQ2
Do you have new questions or want to ask further?	RQ2
Do you think these contents affect your current thoughts or decisions?	RQ3
Did they change your original plans or add new options?	RQ3
Do you think you will refer to these suggestions in reality?	RQ3
Why did the score change? Can you review and explain?	RQ3
Do you think the LLM’s answer caters to you, i.e., deliberately saying nice things?	RQ2

Table 4: Among the detailed demographics of the 20 LLM-fortunetelling interview participants, “LLM usage exp.” indicates how frequently they use large language models in daily life, while “LLM fortunetelling exp.” indicates how frequently they employ LLMs for fortunetelling. Under “Fortunetelling Topics,” the four columns represent different aspects: “In the Past” lists the topics participants mentioned when recalling previous fortunetelling experiences, “Past Type” indicates the type of fortunetelling in participants’ past experiences, “During Experiment” lists the topics they actually used in the study, and “Rating Type” indicates the type of fortunetelling in the experiment. “Belief” means confidence in a predictive outcome (*e.g.*, “Will my paper be accepted?”), while “Decision” means the likelihood of taking a specific action (*e.g.*, “Will I transfer departments?”).

No.	Gender	Age	Occupation	LLM usage exp.	LLM fortunetelling exp.	Fortunetelling Topics				Confidence in Decision	
						In the Past	Past Type	During Experiment	Rating Type	Pre	Post
1	Female	20	Student in Finance	Frequently	Often	Exam prediction	Belief	Career direction	Belief	15	15
2	Male	22	Student in Finance	Often	Often	General fortune prediction	Decision	Romantic fortune prediction	Belief	15	15
3	Male	25	Student in Journalism	Frequently	Often	PhD decision	Decision	Paper acceptance prediction	Belief	75	85
4	Female	20	Student in Sociology	Often	Occasionally	Breakup prediction	Belief	Academic and career planning	Belief	80	80
5	Male	20	Student in Artificial Intelligence	Often	Occasionally	Earthquake probability prediction	Belief	General fortune prediction	Belief	70	80
6	Female	23	Student in Geomatics Engineering	Often	Often	Romantic fortune prediction	Belief	General fortune prediction	Belief	90	95
7	Female	23	Student in Law	Frequently	Frequently	Graduate study prediction	Belief	Love timing prediction	Belief	65	70
8	Female	20	Student in Law	Often	Frequently	Lost item divination	Belief	Academic prospect prediction	Belief	60	80
9	Female	23	Technology company employee	Frequently	Frequently	Resignation time prediction	Belief	Promotion and salary raise prediction	Belief	100	100
10	Female	26	Education industry employee	Frequently	Frequently	Civil service job selection	Decision	Job change decision	Decision	45	50
11	Female	23	Technology company employee	Frequently	Frequently	Work emotion analysis	Decision	Weight loss goal prediction	Belief	80	60
12	Female	39	Self-employed individual	Often	Occasionally	Annual fortune prediction	Belief	Earning potential prediction	Belief	60	90
13	Female	20	Student in Law	Often	Often	General fortune prediction	Belief	Academic prospects prediction	Belief	50	60
14	Male	20	Student in Finance	Frequently	Often	Graduate program selection	Decision	Academic prospects prediction	Belief	80	90
15	Male	31	Technology company employee	Frequently	Occasionally	Civil Service exam prediction	Belief	Career achievement prediction	Belief	50	60
16	Female	43	Self-employed	Often	Occasionally	Investment decision	Decision	Work prospect prediction	Belief	20	40
17	Female	23	State-owned enterprise employee	Frequently	Often	Daily fortune prediction & Travel decision	Decision	Job transition decision	Decision	80	90
18	Female	20	Student in Applied Meteorology	Often	Frequently	Competition participation decision	Decision	Romance fortune prediction	Belief	100	100
19	Female	20	Media sector employee	Occasional	Occasionally	Career path decision	Decision	Career development prediction	Belief	70	82
20	Female	21	Technology company employee	Often	Often	General fortune prediction	Belief	Career continuity decision	Decision	60	50

Table 5: Prompt structures used in LLM-based fortune-telling prompts, identified through a content analysis of 1,045 social media posts.

Categories	Codes	Sample	Posts
Prompt Structure (Google 2025; Liu et al. 2023)	Character	“You are a professional researcher of traditional Chinese numerology.”	137 (13.11%)
	Knowledge	“You are familiar with the Book of Fortune and Poverty, Divine Peak Examination, etc...”	159 (15.22%)
	Rule	“For a male born in a Yang year, arrange in order; for a female born in a Yin year, arrange in order. For reverse order, use the month pillar as the reference.”	120 (11.48%)
	Personal Information	“I was born on [year] [month] [day] at [hour]:[minute], female, in [place of birth].”	219 (20.96%)
	Task	“Please analyze my birth chart as thoroughly as possible.”	228 (21.82%)
	Output Format	“Please provide the output in the following format: 1) Overall fortune (brief summary) 2) Career and studies 3) Love and relationships 4) Wealth and finances...”	49 (4.69%)
	Tone	“Please explain it in simple and easy-to-understand language.”	26 (2.49%)
	Required Explanation	“Please explain the basis for your judgment.”	47 (4.50%)
	None	-	801 (76.65%)