

Soft Information, Hard Decisions: AI Advising^{*}

Jing Huang Shumiao Ouyang

September 2025

Preliminary and References Incomplete

Abstract

We model AI-based advising via large language models (LLMs) by introducing preference uncertainty—capturing soft information—alongside standard fundamental uncertainty, and we validate model predictions through LLM-driven simulations. While human advisors can elicit soft information, their misaligned incentives lead to information loss (Crawford and Sobel, 1982). In contrast, LLMs are unbiased, but digitizing soft information is challenging: it is difficult to formalize and LLMs’ limited memory prevents synthesis across extended conversations. We model this communication as the investor’s optimal stopping problem with Brownian information flow and solve it in closed form. Simulations using multi-round LLM advising confirm our theoretical frictions and are benchmarked against standard portfolio questionnaires. Results show that deeper conversations mainly help investors clarify and learn their own objectives, improving decisions; however, impatience or early stopping reduces recommendation quality.

JEL Classification: G11, G41, D81, O33, C45, C63, D91, A13

Keywords: Artificial Intelligence, Large Language Models, Hardening Soft Information, Financial Advising

^{*}Huang: Mays Business School, Texas A&M University; email: jing.huang@tamu.edu; Ouyang: Saïd Business School, University of Oxford; email: shumiao.ouyang@sbs.ox.ac.uk. For helpful comments, we thank Snehal Banerjee, Bo Bian, Peter Cziraki, Zhiguo He, Dan Luo, Shane Miller, Hongwei Mo, Cecilia Parlato, Jian Sun, Yao Zeng, Leifu Zhang, and seminar participants at Texas A&M University (Mays) and University of Michigan (Ross). Shumiao Ouyang thanks Oxford RAST for their support, particularly Andreas Charisiadis and Shang Wang for their excellent research assistance. All errors are our own.

1 Introduction

The proliferation of artificial intelligence (AI) in advisory roles—spanning financial planning, healthcare, and education—has created a compelling paradox in professional services. While AI advisors demonstrate objective advantages including elimination of conflicts of interest, scalable delivery, and performance comparable to human professionals, their market penetration remains surprisingly limited. In financial markets, the paradox is especially pronounced: although robo-advisors create portfolios with historical returns comparable to those of professional benchmarks (Fieberg et al., 2024) and assist seasoned investors in minimizing behavioral errors (Guo et al., 2022), they account for less than 2% of global assets under management as of 2025 and are expected to experience modest growth in the near future.¹ This limited adoption occurs even as traditional human advisors face well-documented incentive problems—commission-driven recommendations, product pushing, and strategic exploitation of client biases (Bolton et al., 2007; Carlin and Manso, 2011; Chalmers and Reuter, 2020). The persistence of potentially biased human advisors over unbiased AI alternatives suggests that current AI systems cannot replicate crucial aspects of human advisory value.

This missing value lies in the processing of “soft information”—the subjective, often unarticulated preferences, constraints, and goals that clients themselves may not effectively communicate or even fully understand themselves (Liberti and Petersen, 2019). Unlike prediction tasks involving external outcomes such as asset returns, advisory excellence requires understanding the individual client’s internal landscape of uncertainty. A human advisor can navigate this uncertainty through extended dialogue, using probing questions, analogies, and examples to help clients communicate their own preferences. However, large language models (LLMs) face fundamental architectural constraints that create a novel form of information loss distinct from the strategic distortions in human advising. When clients cannot articulate their needs through effective prompts, and when LLMs’ limited memory prevents synthesis across extended conversations, the result is generic recommendations that overlook the nuanced requirements underlying high-stakes decisions.

This paper develops a formal framework for AI-based advising via LLMs, and validates model predictions through LLM-driven simulations that generate multi-turn, role-structured conversations. Alongside standard fundamental uncertainty in financial advising, we introduce “preference uncertainty” to capture soft information, and model the communication of soft information as the investor’s optimal stopping problem with Brownian information flow. The comparison between human versus LLM-based advisors is clear. Human advisors efficiently elicit soft information, but their misaligned incentives lead to strategic misreport-

¹See analysis in <https://resoinsights.com/insight/robo-advisors-in-wealth-management/>

ing and information loss about asset fundamentals as in the standard cheap talk models (Crawford and Sobel, 1982). By contrast, LLM-based advisors are unbiased, yet we identify a novel source of information loss stemming from the frictions inherent in digitizing soft information.

In the model, the investor has a quadratic loss utility function and faces two layers of uncertainty. The first—uncertainty about the fundamental state of the world—is standard in the financial advising literature. The second, which is novel in this paper, is uncertainty about the investor’s own objective, which we call “preference uncertainty,” and it captures soft information. Specifically, the investor does not know which of the fundamental state to match (preference uncertainty); for either potential states, she does not observe its realization (fundamental uncertainty). Hence, the investor herself is “confused”: she does not fully know her preferences *ex ante* and can learn them through consultation with an advisor. For example, while she may know her income status and lifetime goals, she lacks knowledge of personal finance and is unaware how such information is related to her investment objectives.

We include the case of consulting a human advisor for comparison. Human advisors can uncover soft information through interactive dialogue—using probing questions, analogies and examples to help clients communicate their own preferences. In our model, the “preference uncertainty” is eliminated with a human advisor: after an initial consultation, the investor learns her preference and is matched with a specialist tailored to her investment goals. However, the human advisor is biased—e.g., toward generating commission fees—and thus inflates the value of asset fundamentals. As in standard cheap talk models (Crawford and Sobel, 1982), this strategic distortion limits how much credible information the advisor can transmit, resulting in information loss about fundamental values.

In contrast, the LLM advisor is biased, but digitizing soft information is difficult due to two reasons. First, by its nature, digitization inevitably incurs information loss—a challenge that is amplified when the investor herself is uncertain. Second, the limited memory of LLMs prevents synthesis across extended conversations, even though each individual prompt is informative. The Transformer architecture underlying most LLMs scales quadratically with input length—which is the cost of LLMs’ intelligence, but makes long-context processing computationally costly and often impractical. As a result, LLMs typically operate within short, stateless context windows, with knowledge reset after each interaction.

We model the communication with the LLM as the investor’s optimal stopping problem with Brownian information flow. Before seeking a recommendation (stopping), the investor discusses her situation, generating public signals about her underlying preferences that she also comes to learn. The LLM, however, only partially learns about these preferences due to its limited memory. In the baseline case, we assume the LLM has one-shot memory and

updates its belief using only the most recent signal, which, in the continuous-time limit, its belief remains fixed at the pretrained prior. We also consider an extension in which the LLM randomly misses each signal. The investor decides when to stop the conversation and seek recommendation, weighing the informational value of continued learning against information costs.

When the investor stops communication, she receives a recommendation from the LLM. Aware that it reflects the LLM’s belief, she chooses an action based on her own posterior belief, partially correcting for the LLM’s misunderstanding. The investor’s stopping value depends on both the residual preference uncertainty and the residual fundamental uncertainty. Because of communication costs and the LLM’s limited memory, soft information is never fully transmitted. Moreover, although the LLM is unbiased, fundamental uncertainty remains due to unresolved preference uncertainty. Uncertain about the client’s preferences, the LLM generates recommendations for an average client through a black-box process that offers little transparency about the underlying states. This output serves only as a noisy signal of the investor’s preferred fundamental state. Taken together, when consulting an LLM, information loss also arises essentially from the inefficiency in communicating soft information.

We solve the investor’s optimal communication policy in closed form for the baseline model. In equilibrium, the investor follows a threshold strategy: she continues interacting with the LLM while her belief about her preference remains within an intermediate range and stops once she becomes sufficiently confident—either above an upper threshold or below a lower one. This policy reflects a trade-off between the benefit of continued learning about her own objective and information costs.

The model yields rich comparative statics that illuminate when AI advising is more or less effective. When the LLM partially learns from the communication, the investor engages in a deeper communication and chooses more extreme stopping thresholds. We show that the investor communicates less when the cost of interaction is high or when the variance of her preference is small, leading to earlier stopping and greater information loss. The investor communicates longer when she faces substantial uncertainty—when communication adds the most value. However, her resulting value from AI advising is still low. We also find that an LLM model that is trained for more specialized clients induces more learning, while increasing fundamental uncertainty reduces overall payoff but does not affect the learning region. These patterns underscore that the inefficiencies in AI advising stem more from frictions in digitizing soft information than from errors in forecasting hard fundamentals.

We derive a set of testable implications that connect the model’s primitives to observable advising behavior and outcomes. For clarity, we organize hypotheses into two groups: those

we can evaluate in controlled LLM-based simulations and those that require observational field data. In the paper we focus on three simulation-testable hypotheses. H1 posits that the primary value of interacting with an LLM advisor is that investors learn about their own initially uncertain preferences; the LLM’s recommendation may remain generic, but dialogue helps investors reduce “preference uncertainty” and choose portfolios closer to their true objectives. H2 predicts that higher time costs induce earlier termination of the conversation, yielding less tailored portfolios. H3 predicts that giving the LLM persistent access to memory improves recommendation quality and narrows the gap with a human advisor.

We complement the theory with prompt-based LLM simulations that generate multi-turn, role-structured conversations approximating real-world advisory exchanges. Investor “ground truth” comes from the 11-question Investor Questionnaire provided by Vanguard; we simulate $n = 500$ profiles (fixed seed), obtain each profile’s rule-based “optimal” stock/bond allocation, and use these as benchmarks in evaluation. The advisor is implemented with OpenAI’s GPT-5 under strict prompting. We compare a memoryless advisor, which is fed only the most recent Q&A and thus approximates fixed priors due to context-length limits, with a memory-augmented advisor that can use full chat history, and a full-information counterfactual that observes the complete profile. Conversations conclude according to an optimal stopping rule, modeled as a 0.10 exogenous probability of termination per round. The investor assesses the costs and benefits in each round and may choose to end the conversation early, with a maximum of eleven questions permitted.

Our findings, drawn from 2,500 conversations (500 profiles each undergoing 5 interactions), reveal several key patterns. First, supporting H1, it’s clear that the act of interacting itself leads to most of the observed improvements: accuracy rises by 14.1 percentage points even before any specific recommendations, and by 15.7 points afterward. Each additional round of exchange adds about 1.03 points to accuracy, while every extra word contributes roughly 0.017 points. Second, in line with H2, we find that ending conversations prematurely—outside the advisor’s control—reduces accuracy by around 2.62 points. Even when accounting for total rounds, this drop remains substantial at about 0.96 points, indicating a real penalty for impatience. Finally, as H3 predicts, access to memory significantly boosts recommendation quality: the system works best when it has full access to information, performs next best with full memory, and does worst with no memory at all.

Methodologically, this paper introduces an innovative approach to testing economic theory by complementing our analytical framework with prompt-based LLM simulations. This technique allows us to generate dynamic, multi-turn conversations that approximate real-world advisory exchanges, bridging the critical gap between theoretical rigor and empirical realism. Unlike traditional laboratory experiments or analytical models, these simulations

enable the direct observation of complex mechanisms, such as information acquisition, belief updating, and optimal stopping, as they unfold within a controlled yet realistic interactive setting. Although recent literature has begun to explore the potential of using LLMs as economic agents to study behavioral patterns and simulate empirical regularities in human decision-making (Horton, 2023; Ouyang et al., 2024), to the best of our knowledge, we are the first to leverage LLM-driven simulations explicitly as a tool for testing economic theories. This approach opens new possibilities for economic research in domains where communication and context-dependent human behavior play a central role, offering a promising path for validating theoretical predictions.

Literature. This paper contributes to two rising strands of research: Artificial intelligence in financial advising, and the technology trend that transforms soft information into hard data. More broadly, theoretical research on AI technologies has primarily focused on their effects on labor (e.g., Ide and Talamas, 2024), while this paper examines their specific application in providing information within advisory roles.²

AI Advising. We build on classic cheap talk and financial advising models and introduce preference uncertainty to capture soft information. We highlight the key features of AI advisors by comparing them with human advisors, who are misaligned à la the classic cheap talk (Crawford and Sobel, 1982). This seminal framework is later applied to analyst settings where reputational and underwriting incentives drive distorted advice—see Bénabou and Laroque, 1992; Ottaviani and Sørensen, 2006; Rüdiger and Vigier, 2019.³ Our paper explores AI advisors as a new alternative and identifies their limitations: while unbiased, AI systems face challenges in interpreting soft, contextual information.

Our findings contribute to the growing literature on financial technology and AI-driven advisory services, extending the review by Mo and Ouyang (2025). Prior work shows that robo-advisors can improve investor outcomes: D’Acunto et al. (2019) find better diversification and reduced behavioral biases, while Rossi and Utkus (2024) report improved indexing, Sharpe ratios, and lower fees, especially for under-diversified investors. Yet challenges remain. Chak et al. (2022) show that even when robo-advice enhances debt choices, low algorithmic trust can limit uptake. Similarly, Andries et al. (2024) find that human advisors’ effectiveness varies with available information. Recent LLM advances offer new potential: Lu et al. (2023) and Fieberg et al. (2024) show LLMs can generate effective, personalized investment advice. Field evidence from Guo et al. (2022) finds experienced, risk-averse in-

²Theoretical research on AI in finance is still emerging. For example, Chen and Han (2024) show supervised AI intensifies agency conflicts.

³Cheap talk is also applied in corporate-governance settings where boards or proxy advisors offer non-binding recommendations (e.g., Levit and Malenko, 2011; Malenko and Malenko, 2019).

vestors benefit most from conversational AI advisors. Building on this, our paper develops a theoretical model of AI vs. human advising, validated through prompt-based simulations where an LLM engages in iterative dialogue with an investor who chooses when to stop.

This raises the question of whether human experts still offer distinct value in an era of automation. While robo-advisors excel at technical tasks, human advisors provide complementary “soft” services like emotional support and trust-building, which aid prudent decision-making (Linnainmaa et al., 2018; Gennaioli et al., 2015). Human discretion remains valuable: Costello et al. (2020) show it improves outcomes by incorporating private context, and Greig et al. (2024) find that hybrid platforms enhance investor retention and confidence. Similarly, Cao et al. (2024) find that while AI excels in forecasts, humans outperform in tasks needing institutional insight, with combined approaches performing best in uncertain, data-scarce settings. Our findings reinforce that human advisors play a complementary role, especially where nuance, context, and trust matter.

Hardening soft information. Our paper is related to the literature of soft versus hard information, as well as the rise of Big Data and machine learning technologies that transform soft, subjective information into hard, objective data. The literature on soft vs. hard information (e.g., Stein, 2002; Liberti and Petersen, 2019) emphasizes that hard information is verifiable and thus transferable within organizations, while soft information is often non-verifiable and modeled as cheap talk (e.g., Bertomeu and Marinovic, 2016; Corrao, 2023). Recent advances in Big Data and AI have made it possible to digitize even contextual, traditionally soft information. He et al. (2024) examines this technology trend in the context of credit market competition. In our paper, investor communication with the LLM is exactly the process that hardens a fraction of the soft information. However, we emphasize that LLMs face memory constraints that limit their ability to fully process and transform soft information.

The remainder of the paper proceeds as follows. Section 2 introduces the model. Section 3 characterizes equilibrium under human advising and the equilibrium under AI advising separately. Section 4 presents our empirical analysis. Section 5 concludes and discusses practical and policy implications as well as avenues for future research.

2 The Model

A decision maker seeks advice from a better informed advisor—either a human or an AI advisor (large language model, or LLM)—before taking action. To fix ideas, we use the context of financial advising throughout the paper and refer to the decision maker as the “investor.” However, the model applies to a broad range of advisory applications.

2.1 Agents

We first introduce the objectives of all agents, then describe their decision-making given their information. The investor’s dynamic communication with the LLM will be separately introduced in Section 2.2.

2.1.1 Investor.

As in canonical models, we build on the cheap talk framework (Crawford and Sobel, 1982) to model the problem with the human advisor. The key innovation here is modeling the core friction in AI advising as the communication of soft information. To capture this, on top of the standard uncertainty about fundamental states, we introduce a second layer of uncertainty to capture soft information: the investor does not fully understand her needs or the optimization problem, which we refer to as “preference uncertainty.” For example, the investor could be unsure about her risk appetite or how to tailor investments given the number of years until her retirement. Even if she observes the asset fundamentals perfectly (no uncertainty about fundamental states), she may still choose an undesirable portfolio allocation.

To characterize this “preference uncertainty” in the simplest way, we assume that the investor has a quadratic loss utility function and matches her action a with an unobservable fundamental state, but she does not what fundamental state to match—either $\tilde{\theta}_1$ or $\tilde{\theta}_0$. That is, she is unsure whether her objective function is

$$-\mathbb{E}(a - \tilde{\theta}_1)^2,$$

or

$$-\mathbb{E}(a - \tilde{\theta}_0)^2.$$

In addition, $\tilde{\theta}_1$ and $\tilde{\theta}_0$ are independent normal random variables whose realizations are unobservable to the investor:

$$\tilde{\theta}_1 \sim N(\mu_1, \sigma_\epsilon^2), \quad \tilde{\theta}_0 \sim N(\mu_0, \sigma_\epsilon^2),$$

where μ_i is the mean of $\tilde{\theta}_i$ and σ_ϵ^2 is the variance of $\tilde{\theta}_i$ for $i \in \{1, 0\}$. Throughout the paper, we use $\tilde{\theta}_1$ and $\tilde{\theta}_0$ to denote the random variables, and θ_1 (θ_0) to denote a specific realization.

Hence, the investor’s utility function reflects two layers of uncertainty: preference uncertainty and fundamental uncertainty. We introduce $\omega \in \{1, 0\}$ to refer to the investor’s underlying preference: $\omega = 1$ for $\tilde{\theta}_1$ and $\omega = 0$ for $\tilde{\theta}_0$. Fundamental uncertainty means

that, for each preference $\omega \in \{1, 0\}$, the investor does not observe the realization of the corresponding fundamental state θ_i —this is also the standard uncertainty in the literature.

We interpret preference uncertainty as soft information, and refer to the two terms interchangeably in the paper. In the baseline example of financial advising, the investor may know her income status and general risk appetite in daily life but does not know how to translate this soft information into portfolio goals. More broadly, this situation arises when an inexperienced decision maker faces complex, personalized decisions—such as choosing medical insurance for the first time or attempting to self-diagnose a medical conditions. As will be discussed later, this soft information ω can be efficiently elicited only through human interaction. In contrast, a large language model (LLM) cannot fully elicit w , as the investor is unable to articulate soft information in the prompt, and the LLM’s short memory prevents it from summarizing ω from a long, partially relevant communication.

We introduce $p \in [0, 1]$ as the investor’s belief about her preference:

$$p \equiv \mathbb{P}^i(\omega = 1). \quad (1)$$

Her prior belief about ω before consulting an investor is p_0 . Then, the investor’s utility is

$$U(a, p, \theta_1, \theta_0) = -p(a - \theta_1)^2 - (1 - p)(a - \theta_0)^2. \quad (2)$$

We assume that the investor seeks advice from either a human advisor or an LLM before making a decision. The key focus of the paper is on consultation with the LLM and how it compares to the human advisor. While our model sheds light on contexts where each advisor may be better suited, we do not emphasize the investor’s optimal choice of her advisor.

2.1.2 Human advisor

It is widely accepted that soft information can be effectively collected via human interactions (Liberti and Petersen, 2019). In our context, the human advisor can uncover the investor’s preference ω through interactive dialogue—asking follow-up questions, offering analogies, and helping the client clarify her preferences when she doesn’t know how to articulate them.

Formally, when the investor chooses the human advisor, a public signal is generated that perfectly reveals the investor’s preference ω . While neither the investor nor the human advisor knows ω beforehand, the investor does possess related information—such as her income, career status, risk appetite, and spending habit—but does not understand how this information maps to ω . The human advisor, who understands this mapping, elicits the relevant information and thereby uncovers ω . Hence, “preference uncertainty” is eliminated,

and both parties know the target fundamental asset is $\tilde{\theta}_\omega$.

After ω is revealed, the financial advising problem with the human advisor follows classic cheap talk framework (Crawford and Sobel, 1982). The human advisor has conflicted interest measured by $b > 0$, and his utility function is

$$U^h(a, \theta_\omega) = -(a - (\theta_\omega + b))^2. \quad (3)$$

In practice, such bias reflects commission-based incentives, leading the advisor to prefer that the investor purchase more securities. That is, the advisor’s preferred action is $\theta_\omega + b$, which is higher than the fundamental value θ_ω of the financial asset.

The consultation proceeds as follows. The human advisor perfectly observes the realization of the fundamental asset value, θ_ω , and then sends an unverifiable signal m^h (with the superscript “ h ” denoting “human”) to the investor; we specify this signal m^h to be the advisor’s recommended action, and call it “recommendation m^h ” interchangeably. The investor processes the information in the recommendation and chooses an action $a(m^h)$ to maximize her expected payoff as defined in Eq. (2),

$$\max_a \mathbb{E}[U(a, \omega, \theta_\omega) | m^h] = -\mathbb{E}[(a - \tilde{\theta})^2 | m^h], \quad (4)$$

where the investor perfectly knows her preference $p = \omega \in \{0, 1\}$ and matches θ_ω . This action in turn determines the payoffs for both parties.

In this setting, the human advisor strategically communicates only the fundamental information θ_ω , not the investor’s preference ω —that is, misreporting ω to induce a preferred action. This simplification is due to our focus on the AI advisor, with the human advisor serving primarily as a benchmark for comparison. In practice, there are different advisors within the advisory firm who specialize in $\tilde{\theta}_1$ and $\tilde{\theta}_0$ (say $\tilde{\theta}_1$ and $\tilde{\theta}_0$ are very different, such as equity investment versus fixed income.) An initial consultation is typically used to identify the investor’s needs, and assign her with the specialist tailored to her specific situation. Since these advisors are different individuals, there is no strategic communication about ω .

2.1.3 LLM’s recommendation.

Since large language models (LLMs) are designed to minimize prediction errors, we assume that LLM is unbiased in its recommendation.⁴ However, there are still two frictions that lead to inefficient communication of the soft information or the investor’s preference ω .

⁴While developers may train LLMs to maximize user subscriptions, this incentive is orthogonal to misaligned recommendations.

First, by its nature, digitizing soft information inevitably incurs information loss (Liberti and Petersen, 2019)—a challenge that is amplified when the investor herself is “confused.” Besides the loss of important contextual cues when translating the investor’s thoughts into prompts, she lacks understanding of what information is relevant and should be provided—precisely why she is uncertain about her preference ω to begin with.⁵

One may argue that, as long as each prompt is informative, the LLM advisor could eventually gather all relevant aspects of the soft information through past interactions. However, this brings us to our second point: LLMs have a “short memory.” LLMs typically operate within short, stateless context windows, with knowledge reset after each interaction.⁶ This is because the Transformer architecture underlying most LLMs scales quadratically with input length, making long-context processing computationally costly and often impractical. Moreover, extended prompts—such as including the entire chat history—do not function as memory effectively and significantly degrades the LLM’s performance (Liu et al., 2023).⁷

In Section 2.2 below, we model the communication with the LLM as the investor’s optimal stopping problem with Brownian information flow about soft information ω . Now we describe what happens when the investor seeks recommendation from the LLM at any given pair beliefs of herself and the LLM’s—both are endogenously determined by the communication process in Section 2.2. Let $\hat{p} \equiv \mathbb{P}^L(\omega = 1) \in [0, 1]$ denote the LLM’s belief, and $\hat{p}$ is public information: the investor observes the LLM advisor’s pretrained prior $\hat{p}_0$ and its belief update in the communication. However, the investor’s belief p is not observed by the LLM advisor. Since the LLM is unbiased, its payoff U^L is its conjectured investor’s payoff under belief $\hat{p}$:

$$U^L(a, \hat{p}, \theta_1, \theta_0) = U(a, \hat{p}, \theta_1, \theta_0) = -\hat{p}(a - \theta_1)^2 - (1 - \hat{p})(a - \theta_0)^2. \quad (5)$$

The LLM perfectly observes the realization of fundamental states, θ_1 and θ_0 . Then, when the investor seeks recommendation for her action, the LLM sends

$$m^L = \hat{p}\theta_1 + (1 - \hat{p})\theta_0. \quad (6)$$

The investor understands that m^L is determined in Eq. (6) under the LLM’s belief $\hat{p}$, but does not observe the fundamental realizations θ_1 and θ_0 . She then chooses her optimal action

⁵While LLMs have made notable progress in initiating conversations that resemble human dialogue, their performance crucially relies on the quality of the input (“garbage in, garbage out.”)

⁶Each response is newly generated based solely on the current context, without incorporating any algorithmic memory of previous responses. Depending on the LLM, systems like ChatGPT may include prior texts from the same chat session in the current context, creating the impression of memory.

⁷While emerging solutions like lightweight key-value caches, vector-store retrieval, and weight-editing offer some mitigation, these approaches are still in development.

$a(m^L, \hat{p})$ to maximize her expected utility, under her belief that $\omega = 1$ with probability p :

$$g(p, \hat{p}|m^L) \equiv \max_a \mathbb{E}[U(a, p, \theta_1, \theta_0)|m^L, \hat{p}] = -p\mathbb{E}[(a - \theta_1)^2|m^L, \hat{p}] - (1 - p)\mathbb{E}[(a - \theta_0)^2|m^L, \hat{p}]. \quad (7)$$

Let $a^*(m^L, \hat{p})$ denote the optimal action and we denote the resulting value as $g(p, \hat{p}|m^L)$. Then, given any pair of the investor's belief p and the LLM's belief $\hat{p}$, the investor's expected utility when she seeks recommendation is:

$$g(p, \hat{p}) \equiv \mathbb{E}[g(p, \hat{p}|m^L)] = \mathbb{E}[U(a^*(m^L, \hat{p}), p, \theta_1, \theta_0)]. \quad (8)$$

As will be introduced later, the investor seeks recommendation when the communication stops. Hence, Eq. (8) captures the investor's stopping value at any pair of beliefs p and $\hat{p}$ in the communication to be introduced next in Section 2.2.

Remark 1. The message space of the LLM is restricted to be a single recommendation m^L . When $\hat{p} \in (0, 1)$, fundamental uncertainty remains regarding the realization of $\tilde{\theta}_\omega$. We argue that the LLM cannot return both underlying fundamentals, θ_1 and θ_0 , which aligns with the black-box nature of AI outputs. Note that our simple two-state structure serves as an abstraction of the complex underlying algorithms. Even if the LLM were to provide intermediate reasoning steps, a typical customer would be unable to infer θ_1 and θ_0 in a meaningful way.

2.2 Communication with the LLM

The investor may initiate multiple rounds of conversations with the LLM to discuss about her needs, which gradually reveals information about ω . The investor and the LLM holds a prior belief of p_0 and $\hat{p}_0$ that $\omega = 1$. The communication is in continuous time, starting at $t = 0$ with an infinite horizon, and each round of chat lasts time dt . The investor is impatient: with a Poisson shock of intensity λ , the communication ends and she receives a recommendation immediately.

At every time t , the investor incurs a cost of $cdt > 0$ if she continues the communication during the interval $[t, t + dt)$. Whenever the communication stops—either by the investor voluntarily or by an exogenous Poisson event, the investor seeks recommendation from the LLM and the game stops. Suppose at time t , the game has not yet stopped exogenously, and investor stops the communication to seek recommendation. The payoff to the investor is then

$$-ct + g(p_t, \hat{p}_t).$$

In addition, signals about the investor's preference ω is gradually revealed to both parties—as we discuss below.

Signals about the soft information ω . Signals about the principal's true preference ω ($\omega = 1$ for matching $\tilde{\theta}_1$ or $\omega = 0$ for $\tilde{\theta}_0$) is revealed by a sequence of signals modeled as a Brownian diffusion process with drift ω . Specifically, the signal process evolves according to

$$ds_t = \omega dt + \sigma dB_t, \quad (9)$$

where $B = \{B_t, \mathcal{F}_t, 0 \leq t \leq \infty\}$ is standard Brownian motion on the canonical probability space. The signal ds_t in (9) is more informative if σ is small. At each time t , the entire history of signals, $\{s_\tau, 0 \leq \tau \leq t\}$, is publicly observable. However, as will be discussed later, the LLM learns from the signals partially due to its short memory.

2.2.1 Belief updating

Investor's belief p . At every time t , the investor's belief that $\omega = 1$, p_t , is conditioned on the history of past communication, or the filtration $\mathcal{F}_t^s$ generated from $\{s_\tau, 0 \leq \tau \leq t\}$. Let f_t^ω denote the density of s_t conditional on ω , which is normally distributed with mean ωt and variance $\sigma^2 t$, i.e., $s_t \sim N(\omega t, \sigma^2 t)$. The investor's posterior belief p_t satisfies the Bayes rule

$$p_t = \frac{p_0 f_t^1(s_t)}{p_0 f_t^1(s_t) + (1 - p_0) f_t^0(s_t)}. \quad (10)$$

The log-likelihood ratio $z_t \equiv \frac{p_t}{1-p_t}$ evolves according to a simple process, with which we derive the evolution of the investor's belief process p . Taking the log-likelihood ratio of (10):

$$z_t = z_0 + \ln \frac{f^1(s_t)}{f^0(s_t)} = z_0 + \frac{1}{\sigma^2} \left(s_t - \frac{t}{2} \right), \quad (11)$$

and thus

$$dz_t = \frac{1}{\sigma^2} \left(ds_t - \frac{1}{2} dt \right).$$

From the investor's perspective, $\omega = 1$ with probability p_t and signal s_t is released according to $ds_t = p_t dt + \sigma dB_t$, where $p_t = \frac{e^{z_t}}{1+e^{z_t}}$. The evolution of investor's belief measured in log-likelihood ratio is then

$$dz_t = \frac{1}{\sigma^2} \left(\frac{e^{z_t}}{1+e^{z_t}} - \frac{1}{2} \right) dt + \frac{1}{\sigma} dB_t. \quad (12)$$

Using Ito’s Lemma and $z_t(p_t) = \ln \frac{p_t}{1-p_t}$, we show that p_t evolves according to the standard binary state Brownian signal formula (for details, see Appendix A.1),

$$dp_t = \frac{p_t(1-p_t)}{\sigma} dB_t. \quad (13)$$

There are a few things worth noting. First, the belief process $\{p_t\}$ is a martingale without a drift term, reflecting prior consistency—the expectation of posterior beliefs equals the prior. Second, the term $\frac{1}{\sigma}$ is the signal-to-noise ratio: the numerator equals the gap in drifts when $\omega = 1$ versus when $\omega = 0$. Finally, the belief p_t is absorbing at $p_t = 0$ or 1 : once the investor becomes certain about her type, there is nothing further to learn.

LLM’s belief $\hat{p}$: baseline. As discussed in Section 2.1.3, LLMs have a short memory. In fact, their knowledge is reset after responding to each prompt, and they operate within short, stateless context windows. Consistent with this technology constraint, our baseline model assumes that the LLM advisor exhibits “one-shot memory”—that is, it updates its belief solely based on signal ds_{t-} from the most recent round of communication. In Appendix A.1, we use a standard Binomial approximation to show that, under the continuous-time setting, this form of memory implies the LLM effectively never updates its belief beyond its prior $\hat{p}_0$:

$$\hat{p}_t = \mathbb{E}[\omega = 1 | ds_{t-}] = \hat{p}_0. \quad (14)$$

Therefore, in the baseline, the LLM’s belief $\hat{p}_t$ remains constant at $\hat{p}_0$, which we interpret as its pre-trained belief—that is, the fraction of $\omega = 1$ in the LLM’s customer base.

This baseline setting is not an oversimplification. Although the LLM’s belief remains constant, it still meaningfully affects model implications because the pre-training can vary. Moreover, in a discrete-time setting, the LLM can update its belief based on the most recent signal, and the key economic insight remains qualitatively consistent with the continuous-time model presented here.

The baseline captures the following scenario in practice: users who are uncertain about their needs and unable to articulate such soft information provide prompts that carry minimal information to the LLM, which, given the last prompt, would then generate recommendations based on a “typical” customer profile. Meanwhile, the user herself gradually learns about her preference over time through the communication, as reflected in the evolution of p_t .

LLM’s belief $\hat{p}$: extension. We also consider the extension where the LLM can update its belief based on the entire communication history, but only *partially* learns from the signals $\{s_t\}$. This case speaks to the future developments, as advances in AI could make the

technology constraint less binding (see Footnote 7 for emerging solutions.) Even under the current technology, one straightforward way to augment the LLM’s memory is to feed the entire chat history into subsequent prompts. Indeed, some LLMs internally incorporate previous interactions from the same chat session into the current context window—for example, ChatGPT does so within a single chat.

However, the LLM’s performance degrades significantly when operating over a long context. The Transformer architecture underlying most LLMs correlates every input token (in the context window) with the whole universe of tokens, making long inputs computationally intensive and less efficient. Empirical work have shown that, LLMs perform best when key information appears at the beginning or end of the input (due to primacy and recency effects), but performance degrades significantly when crucial information is placed in the middle (Liu et al., 2023).

In the model, we assume that for signal ds_t in each round of interaction, with probability $\kappa \in [0, 1)$, the LLM “absorbs” this signal and uses it to update its belief $\hat{p}_t$; with probability $1 - \kappa$, the LLM misses this signal and does not update its belief. We assume that the events of missing the signal are independent across time.

We use the LLM’s log-likelihood $\hat{z}_t = \ln \frac{\hat{p}_t}{1-\hat{p}_t}$ to discuss the evolvment of its belief. Importantly, if the LLM misses the signal s_t , it is equivalent to case where it receives an alternative signal $\tilde{s}_t$ whose noise is infinite. In this case, the change in $\hat{z}_t$, which takes the same form of Eq. (12), is

$$d\hat{z}_t = \lim_{\tilde{\sigma} \rightarrow \infty} \frac{1}{\tilde{\sigma}^2} \left(d\tilde{s}_t - \frac{1}{2} dt \right) = 0.$$

In the other case, if the LLM absorbs this signal, the update in its belief should be the same as that of the investor; that is, $d\hat{z}_t = dz_t$.

Therefore, when measured in log-likelihood ratio, the LLM’s belief update is exactly κ fraction of that of the investor,⁸

$$d\hat{z}_t = \kappa dz_t. \tag{15}$$

Or equivalently,

$$\hat{z}_t = \hat{z}_0 + \int_0^t d\hat{z}_s = \hat{z}_0 + \kappa(z_t - z_0). \tag{16}$$

Eq. (16) implies that the LLM’s belief is a function of the investor’s belief, $\hat{p}_t(p_t)$.

⁸One can consider the following microfoundation. Let $\{X_i\}$ be iid Bernoulli random variables with $\mathbb{P}(X_i = 1) = \kappa$. Then

$$\hat{z}_t - \hat{z}_0 = \int_0^t \mathbf{1}_{X_s} dz_t = \lim_{n \rightarrow \infty} \sum_{s=0}^{tn-1} \left[\lim_{K \rightarrow \infty} \sum_{k=1}^K \mathbf{1}_{X_{sk}} \frac{1}{K} \left(z_{\frac{s+1}{n}} - z_{\frac{s}{n}} \right) \right] = \lim_{n \rightarrow \infty} \sum_{s=0}^{tn-1} \kappa \left(z_{\frac{s+1}{n}} - z_{\frac{s}{n}} \right) = \kappa \int_0^t dz_t,$$

where the second last equation applied the Law of Large Numbers.

Note that this extension nests the baseline model when $\kappa = 0$: the LLM never learns from the conversation, and its belief is fixed at prior q_0 . In the other extreme, $\kappa = 1$, the LLM shares the same belief as the principal. In this extension, we assume that κ is sufficiently small, so that the equilibrium shares the same as in the baseline model.

Remark 2. In our model, the investor’s belief p_t is a sufficient statistics of past communication to her. However, she cannot summarize this belief into a prompt to elicit a recommendation from the LLM. This limitation reflects the nature of soft information: while the investor understands p_t in her thoughts, she is unable to articulate it clearly in her prompt. Moreover, if she were able to partially convey p_t to the LLM, the setting would closely resemble our extension in which the LLM partially learns from the entire communication history.

Remark 3. In our model, the investor is rational but the LLM in our model is not. Given the LLM’s learning in (16), the investor conjectures the LLM’s posterior belief $\hat{p}_t(p_t)$ based on her belief p_t . If the LLM were rational, it would back out the investor’s belief p_t from $\hat{p}_t$ as well. If both parties were rational, asymmetric information would not arise in our extension of the LLM’s partial learning.

2.3 Equilibrium Definition

Consultation with the human advisor. Suppose the investor consults with the human advisor and her preference is revealed as ω .

Definition 1. When the advisor is human, an equilibrium consists of the advisor’s signaling rules, denoted by $\pi(m^h|\theta_\omega)$, and an action rule for the investor $a(m^h)$, such that:

- (i) for any realization of θ_ω , $\int \pi(m^h|\theta_\omega)dm^h = 1$, and if m^{h*} is on the support of $\pi(\cdot|\theta_\omega)$, then m^* solves $\max_{m^h} U^h(a(m^h), \theta_\omega)$ where $U^h(a, \theta_\omega)$ is given in Eq. (3).
- (ii) for each m^h , $a(m^h)$ solves $\max_a \mathbb{E}[U(a, \omega, \theta_\omega)]$ as in Eq. (4).

Consultation with the LLM. Given the evolution of her belief $\{p_t\}$, the investor faces an optimal stopping problem

$$(SP) \quad \sup_{\tau \geq 0} \mathbb{E}^\omega \left\{ \int_0^\tau e^{-\lambda t} [-c + \lambda g(p_t, \hat{p}_t(p_t))] dt + e^{-\lambda \tau} g(p_\tau, \hat{p}_\tau(p_\tau)) \right\},$$

where her belief p_t evolves according to Eq. (13). With probability $e^{-\lambda \tau}$, the game has not stopped by time τ , allowing the investor to end communication voluntarily and receive

$g(p_\tau, \hat{p}_\tau)$ in Eq. (8). In the baseline case, $\hat{p}_\tau = \hat{p}_0$, while in the extension, $\hat{p}_\tau(p_\tau)$ is a function of p_τ as implied by Eq. (16). At any earlier time $t < \tau$, the probability that the game has not yet ended is $e^{-\lambda t}$; during the next time interval dt , the investor incurs communication cost $c dt$, and with probability λdt , the game ends exogenously and she receives $g(p_t, \hat{p}_t)$.

Definition 2. When the investor consults the LLM, an equilibrium consists of the investor's stopping rule τ , the LLM's recommendation $m^L(\hat{p})$ and the investor's action $a(m^L)$ such that:

- (i) Stopping time τ solves the investor's communication problem (*SP*);
- (ii) The process of p_t is given by (12), and the process of $\hat{p}_t$ is implicitly given by (16), with the baseline case of $\kappa = 0$ and $\hat{p}_t = \hat{p}_0$;
- (iii) Given any belief $\hat{p} \in [0, 1]$, the LLM's recommendation $m^L(\hat{p})$ solves $\max_{m^L} U^L(m^L, \hat{p}, \theta_1, \theta_0)$ and satisfies Eq. (6);
- (iv) Given any pair of beliefs $p \in [0, 1]$ and $\hat{p} \in [0, 1]$, the investor's action $a(m^L)$ solves $\max_a \mathbb{E}[U(a, p, \theta_1, \theta_0) | m^L(\hat{p})]$.

3 Equilibrium

We briefly characterize the equilibrium under the human advisor in Section 3.1 as a benchmark for comparison with AI advising. In Section 3.2, we construct the equilibrium with the LLM advisor and solve it in closed form.

3.1 Human advisor: cheap talk

When consulting a human advisor, communication about the soft information is efficient and ω becomes public. However, as in the standard cheap talk literature (Crawford and Sobel, 1982), misaligned objectives between the investor and advisor lead to information loss about the fundamental state θ_ω . Since the human advisor is biased with $b > 0$ and his message m^h is non-verifiable, he has an incentive to exaggerate the fundamental state θ_ω by sending a higher message m^h to induce a higher action a . Anticipating this, the investor rationally discounts the overly optimistic messages. In equilibrium, only partial information about $\tilde{\theta}$ can be credibly communicated. The following proposition, as a direct application of Theorem 1 of Crawford and Sobel (1982), characterizes the equilibrium.

Proposition 1. *Suppose b is sufficiently small. There exists a positive integer $N(b)$ such that for every N with $1 \leq N \leq N(b)$, there exists at least one equilibrium where the support*

of θ_ω is partitioned into intervals by $-\infty = \theta_0 < \theta_1 < \dots < \theta_N = \infty$, and the human advisor sends a distinct message m_i^h for each interval $[\theta_{i-1}, \theta_i)$. At each threshold θ_i , the advisor is indifferent between adjacent messages:

$$\theta_i = \frac{a(m_i^h) + a(m_{i+1}^h)}{2} - b. \quad (17)$$

Given message m_i^h , the investor takes action

$$a_i = \mathbb{E}[\theta \mid \theta \in [\theta_{i-1}, \theta_i)] = \mu_\omega + \sigma_\epsilon \cdot \frac{\phi(\theta_{i-1}) - \phi(\theta_i)}{\Phi(\theta_i) - \Phi(\theta_{i-1})}, \quad (18)$$

where ϕ and Φ are respectively the PDF and CDF of fundamental state $\theta_\omega \sim N(\mu_\omega, \sigma_\epsilon^2)$.

Proposition 1 shows that the equilibrium takes the form of a partitional signal structure. For each distinct message m_i corresponding to an interval $[\theta_{i-1}, \theta_i)$, the investor's optimal action in Eq. (18) is the conditional expectation—the truncated normal mean over the interval, reflecting her quadratic loss utility in (2). In equilibrium, the advisor must be indifferent at each threshold θ_i , which requires $U^h(a(m_i^h), \theta_i) = U^h(a(m_{i+1}^h), \theta_i)$, or equivalently (17).

We focus on the equilibrium with the most informative signal structure (i.e., the highest number of partitions N). Since preference uncertainty is resolved, the investor's expected utility is depends solely on the residual fundamental uncertainty about θ_ω :

$$\mathbb{E}[U(a(m^h), \omega, \theta_\omega)] \equiv - \sum_{i=1}^N \mathbb{P}([\theta_{i-1}, \theta_i)) \cdot \text{Var}(\theta_\omega \mid [\theta_{i-1}, \theta_i)),$$

where $\text{Var}(\theta_\omega \mid [\theta_{i-1}, \theta_i))$ is the variance of a truncated normal distribution over $[\theta_{i-1}, \theta_i]$.⁹ The investor enjoys a higher payoff if more information about θ_ω is transmitted, and the extent of information loss is determined by the magnitude of the advisor's bias b . As b increases, the number of partitions decreases, reducing the investor's payoff. In the limit as $b \rightarrow \infty$, communication becomes uninformative (babbling equilibrium), and the principal learns only her preference.

⁹The explicit expression is

$$\text{Var}(\theta_\omega \mid [\theta_{i-1}, \theta_i)) = \sigma_\epsilon^2 \left[1 + \frac{\frac{\theta_{i-1} - \mu_\omega}{\sigma_\epsilon} \cdot \phi(\theta_{i-1}) - \frac{\theta_i - \mu_\omega}{\sigma_\epsilon} \cdot \phi(\theta_i)}{\Phi(\theta_i) - \Phi(\theta_{i-1})} - \left(\frac{\phi(\theta_{i-1}) - \phi(\theta_i)}{\Phi(\theta_i) - \Phi(\theta_{i-1})} \right)^2 \right],$$

where ϕ and Φ denote the PDF and CDF of $\tilde{\theta}_\omega \sim N(\mu_\omega, \sigma_\epsilon^2)$, respectively.

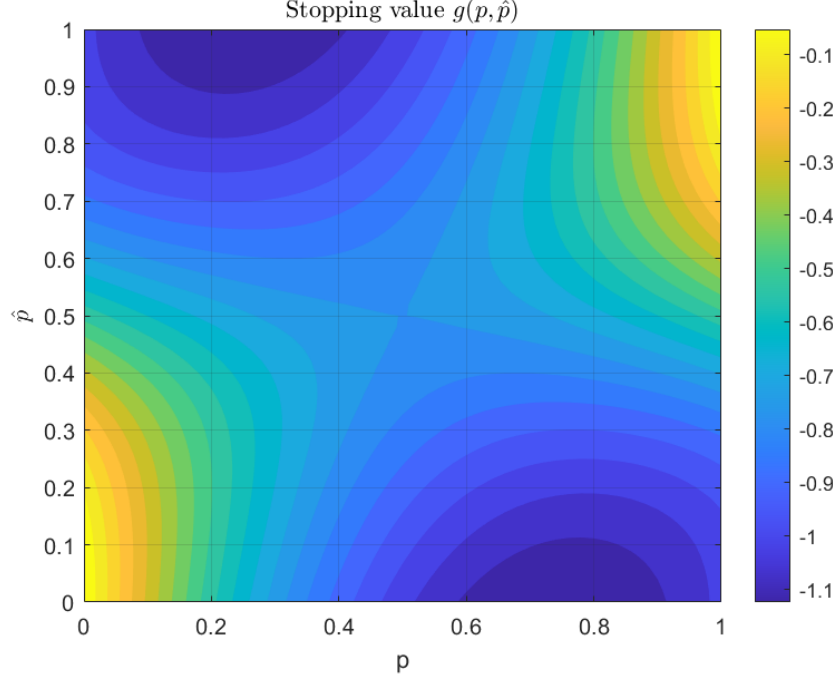


Figure 1: **Investor's stopping value** $g(p, \hat{p})$. Figure shows the investor's expected utility when seeking recommendation, $g(p, \hat{p})$, as a function of the investor's belief p in the horizontal axis and the LLM's belief $\hat{p}$ in the vertical axis. Parameters: $\mu_1 = 1$, $\mu_0 = 0$, $\sigma_\epsilon = 1$.

3.2 AI advising: equilibrium construction

We begin by analyzing the investor's expected payoff given any pair of beliefs $(p, \hat{p})$, highlighting the effect of inefficient communication of the soft information ω . We then construct the equilibrium of interest and provide a closed-form solution.

3.2.1 Stopping value.

Recall that the investor's expected payoff when seeking recommendation is given in (8) in Section 2.1.3. The following lemma characterizes its value.

Lemma 1. (*Stopping value*) *Given any pair of the investor's belief p and the LLM's belief $\hat{p}$ about $\omega = 1$, the investor's expected payoff when seeking recommendation is*

$$g(p, \hat{p}) = -\sigma_\epsilon^2 \frac{p(1 - \hat{p})^2 + (1 - p)\hat{p}^2}{\hat{p}^2 + (1 - \hat{p})^2} - p(1 - p) \left[(\mu_1 - \mu_0)^2 + \frac{(2\hat{p} - 1)^2 \sigma_\epsilon^2}{\hat{p}^2 + (1 - \hat{p})^2} \right]. \quad (19)$$

Figure 1 provides an illustration of $g(p, \hat{p})$. The core friction of AI advising arises from the residual preference uncertainty—inefficient communication of the soft information ω . When both the investor and the LLM are certain and agree on ω —that is, $p = \hat{p} = 1$ or $p = \hat{p} = 0$ —the investor enjoys the highest possible payoff, $g(p, \hat{p}) = 0$.

Otherwise, we rewrite (19) as follows to clarify inefficiency affects the investor's payoff:

$$g(p, \hat{p}) = - \underbrace{\left[p \text{Var}(\tilde{\theta}_1 | m^L) + (1 - p) \text{Var}(\tilde{\theta}_2 | m^L) \right]}_{MSE} - p(1 - p) \mathbb{E}[\Delta \mu_{\theta|m^L}^2 | m^L], \quad (20)$$

where $\Delta \mu_{\theta|m^L} = \mathbb{E}[\tilde{\theta}_1 | m^L] - \mathbb{E}[\tilde{\theta}_0 | m^L]$ is the gap in the conditional mean. The last term captures the direct cost of preference uncertainty. This cost is small when the distributions of $\tilde{\theta}_1$ and $\tilde{\theta}_0$ are close—so that distinguishing between them is unnecessary (i.e., $\Delta \mu_{\theta|m^L}$ is small), or when the investor herself is certain about her preference (i.e., p is close to 1 or 0) and can adjust her action accordingly towards $\tilde{\theta}_\omega$.

In addition, even though the LLM is unbiased, the investor is also subject to residual fundamental uncertainty that arises from the preference uncertainty, as captured by the first two terms in (20). Since the LLM's recommendation m^L in (6) is a weighted average of potential fundamentals, it is a noisy signal for either fundamental θ_1 or θ_0 . (Remark 1 gives a short discussion why the LLM does not send (θ_1, θ_0) as message directly.)

As illustrated in Figure 1, the investor's payoff is higher when she is more certain about her preference, that is, as p approaches 0 or 1. This creates an incentive to continue communication about ω before ultimately seeking a recommendation. In the dynamic communication problem analyzed next, the investor trades off the benefit of learning more about ω against the ongoing cost of communication.

In addition, the principal's payoff is higher when the LLM's belief is more aligned with hers—along the 45-degree line in Figure 1. In the baseline case, where the LLM has only one-shot memory based on ds_{t-} so its belief remains fixed at its pre-trained value $\hat{p}_t = \hat{p}_0$, the investor's payoff is high when she is the typical customer in the LLM's pre-training. In the extension case, where the LLM partially learns from the communication, the investor would enjoy a higher payoff as alignment improves over time.

Note that $g(p, \hat{p})$ is the stopping value in the dynamic communication. Since $\hat{p} = \hat{p}_0$ in the baseline and $\hat{p}$ is a function of p in the extension, we can express the stopping value as a function of the investor's belief: $g(p) \equiv g(p, \hat{p}(p))$.

3.2.2 Value function and optimal policy

The dynamic communication problem is stationary and the value function of the investor depends only on her belief p . Let $V(p)$ denote the investor's value in state p . We conjecture that there exists $0 \leq \underline{p} < \bar{p} \leq 1$ such that the investor continues the communication with the LLM if $p \in (\underline{p}, \bar{p})$.

Outside of the continuation region, $p \notin (\underline{p}, \bar{p})$, the investor immediately seeks recommen-

ation from the LLM and the game stops. The value function is simply the stopping value $g(p) \equiv g(p, \hat{p}(p))$ in Lemma 1.

For all $p \in (\underline{p}, \bar{p})$, the investor incurs cost c to initiate another round of interaction with the LLM. With probability λdt , the communication ends exogenously and she receives the recommendation immediately. Otherwise, the investor receives her continuation payoff. Hence,

$$V(p) = -cdt + \lambda dt g(p) + (1 - \lambda dt) \mathbb{E}_p[V(p + dp)]. \quad (21)$$

Applying the Ito's formula to the right-hand side of (21) gives

$$V(p) \approx -cdt + \lambda dt g(p) + (1 - \lambda dt) \mathbb{E}_p[V(p) + V'(p)dp + \frac{1}{2}V''(p)(dp)^2].$$

Using the law of motion is given in (13), and taking the limit of $dt \rightarrow 0$ gives a linear second-order differential equation for the investor's value function in the continuation region:

$$-c + \lambda(g(p) - V(p)) + \frac{p^2(1-p)^2}{2\sigma^2}V''(p) = 0. \quad (22)$$

All solutions of the differential equation take the following form:

$$V(p) = Q(p) + C_1 p^{\frac{1}{2}+\gamma}(1-p)^{\frac{1}{2}-\gamma} + C_2 p^{\frac{1}{2}-\gamma}(1-p)^{\frac{1}{2}+\gamma}, \quad (23)$$

where $\gamma \equiv \sqrt{2\sigma^2\lambda + \frac{1}{4}}$. The term $Q(p)$ is one particular solution to (22), and under the baseline case where $\hat{p} = \hat{p}_0$, we provide the closed-form $Q(p)$ in Appendix A.3. The two constants C_1 and C_2 are yet to be determined.

The constants are pinned down by the value-matching conditions,

$$V(\underline{p}) = g(\underline{p}), \quad (24)$$

$$V(\bar{p}) = g(\bar{p}). \quad (25)$$

Finally, the two equilibrium belief thresholds $\underline{p}$ and $\bar{p}$ satisfy two smooth pasting conditions, which are required such that the investor's strategy solves (SP) . Specifically,

$$V'(\underline{p}) = g'(\underline{p}), \quad (26)$$

$$V'(\bar{p}) = g'(\bar{p}). \quad (27)$$

Baseline case: $\kappa = 0$. In the baseline case, the LLM has only one-shot memory and makes its recommendation based solely on ds_{t-} , and its belief remains fixed under the dynamic setting: $\hat{p}_t = \mathbb{E}[\omega = 1 | ds_{t-}] = \hat{p}_0$.

In this case, we can show that the investor's optimal stopping thresholds $\underline{p}$ and $\bar{p}$ are symmetric with respect to $\frac{1}{2}$. To see this, $v(p) = V(p) - g(p)$ captures the option value of waiting. The differential equation (22) could be rewritten as

$$\lambda v(p) = \frac{p^2(1-p)^2}{2\sigma^2} g''(p) - c + \frac{p^2(1-p)^2}{2\sigma^2} v''(p).$$

Importantly, $g(p)$ is a quadratic function of p so $g''(p)$ is a constant. It is clear to see that both the flow benefit of waiting, $\frac{p^2(1-p)^2}{2\sigma^2} g''(p) - c$, and the volatility $\frac{p^2(1-p)^2}{2\sigma^2}$ are symmetric for p around $\frac{1}{2}$. Intuitively, when the LLM's belief is fixed, the labels $\omega = 1$ or 0 for the investor's preference type are interchangeable. In contrast, in the extension case where $\kappa > 0$, this symmetry breaks down: the flow benefit of waiting depends on how the LLM's belief evolves, which in turn depends on its prior.

Remark 4. The quadratic form of $g(p)$ arises in many other applications as well. When an economic agent's payoff is linear in her information, as she also chooses an action based on that information, the resulting value function becomes quadratic in the information.

The following proposition summarizes the equilibrium in the baseline case.

Proposition 2. *When $p \notin (\underline{p}, \bar{p})$, the investor stops communication and receives $g(p, \hat{p}_0)$ in (19). When $p \in (\underline{p}, \bar{p})$, the investor's value is*

$$V(p) = Q_0(p) + C \left[p^{\frac{1}{2}+\gamma}(1-p)^{\frac{1}{2}-\gamma} + p^{\frac{1}{2}-\gamma}(1-p)^{\frac{1}{2}+\gamma} \right], \quad (28)$$

where $\gamma = \sqrt{2\sigma^2\lambda + \frac{1}{4}}$, and $Q_0(p)$ is a particular solution given in (39) in Appendix A.3. The constant C is determined by the value matching condition $V(\bar{p}) = g(\bar{p})$. The optimal thresholds $\underline{p}, \bar{p}$ are symmetric: $\bar{p} + \underline{p} = 1$, and $\bar{p}$ is determined by smooth pasting $V'(\bar{p}) = g'(\bar{p})$.

Figure 2 illustrates the equilibrium. There are a few things worth noting. Since the LLM's belief is fixed at $\hat{p}_t = \hat{p}_0 > \frac{1}{2}$, the investor's stopping value $g(p)$ and value function $V(p)$ are tilted upwards for $p > \frac{1}{2}$, where her belief p is more aligned with the LLM's. However, the optimal stopping thresholds $\underline{p}, \bar{p}$ are symmetric around $p = \frac{1}{2}$. In the shaded continuation region, the investor's value function $V(p)$ lies above the immediate stopping payoff $g(p)$, so she chooses to continue communication. Notably, even if the preference uncertainty is fully resolved, that is $p = 1$ or $p = 0$, the investor's payoff $U < 0$ due to the residual fundamental uncertainty: since the LLM does not learn, its recommendation m^L is a noisy signal of the true fundamental θ_ω .

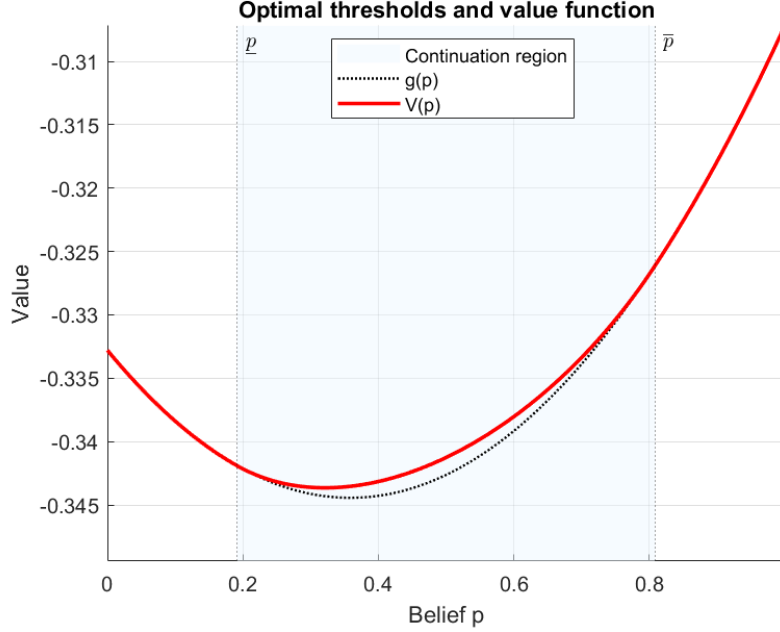


Figure 2: **Optimal policy and value function (baseline $\kappa = 0$).** The figure shows the continuation region $p \in (\underline{p}, \bar{p})$ (shaded area), the investor's value function $V(p)$ (red line) and stopping value $g(p) \equiv g(p, \hat{p}_0)$ (black dotted). Parameters: $\mu_1 = 0.3$, $\mu_0 = 0$, $\sigma_\epsilon = 0.8$, $\sigma = 0.3$, $\lambda = 0.3$, $c = 0.045$, $\hat{p}_0 = 0.51$.

LLM's partial learning: the case of $\kappa > 0$. In this case, for each signal s_t generated in the communication, the LLM absorbs the signal with probability $\kappa > 0$. The baseline case where the LLM only has a one-shot memory could be nested as $\kappa = 0$. For this extension, we consider a sufficiently small κ , under which the structure of the continuation region in Proposition 2 remains robust.

As discussed in Section 2.2.1, the update in the LLM's log-likelihood ratio $\hat{z}_t \equiv \ln \frac{q_t}{1-q_t}$ is κ fraction of that of the investor. The following lemma summarizes the LLM's belief process.

Lemma 2. *The log-likelihood ratio of the LLM's belief $\hat{z}_t \equiv \ln \frac{q_t}{1-q_t}$ satisfies*

$$\hat{z}_t = \hat{z}_0 + \kappa(z_t - z_0).$$

Accordingly, the LLM's belief is a function of the investor's belief p_t :

$$\hat{p}_t(p_t) = \frac{\frac{\hat{p}_0}{1-\hat{p}_0} \left[\frac{(p_t/p_0)}{(1-p_t)/(1-p_0)} \right]^\kappa}{1 + \frac{\hat{p}_0}{1-\hat{p}_0} \left[\frac{(p_t/p_0)}{(1-p_t)/(1-p_0)} \right]^\kappa}.$$

Under a sufficiently small κ , the investor continues communication when $p \in (\underline{p}, \bar{p})$ and stops immediately otherwise. Since the LLM also partially learns from the communication, the optimal thresholds $\underline{p}$ and $\bar{p}$ are no longer symmetric around $\frac{1}{2}$ when $\hat{p}_0 \neq \frac{1}{2}$. Instead,

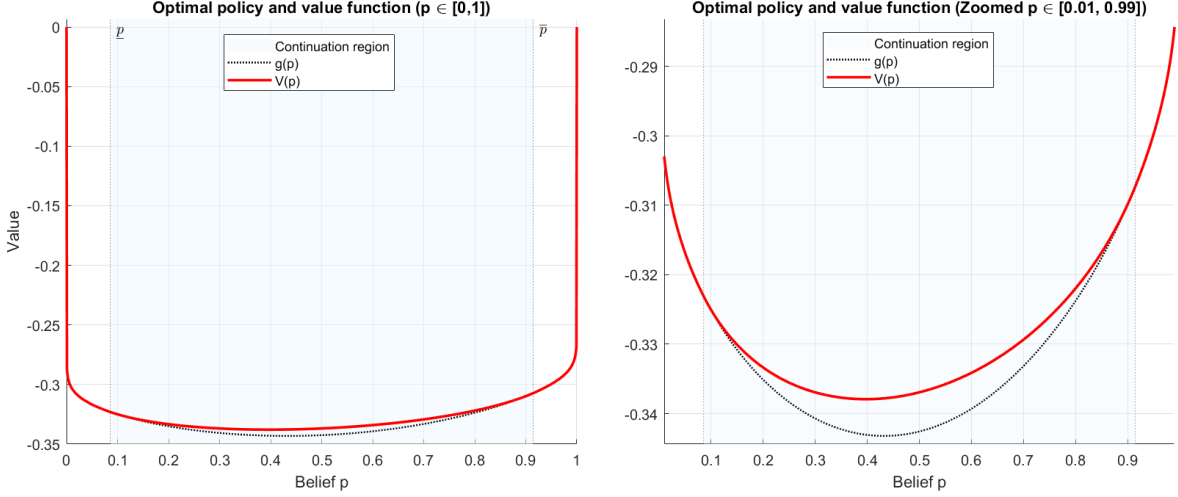


Figure 3: **Optimal policy and value function (extension $\kappa > 0$).** The left panel shows the full plot for $p \in [0, 1]$ and the right panel shows the zoomed plot $p \in (0.01, 0.99)$. The figure shows the continuation region $p \in (\underline{p}, \bar{p})$ (shaded area), the investor's value function $V(p)$ (red line) and stopping value $g(p) \equiv g(p, \hat{p}_0)$ (black dotted). Parameters: $\mu_1 = 0.3$, $\mu_0 = 0$, $\sigma_\epsilon = 0.8$, $\sigma = 0.3$, $\lambda = 0.4$, $c = 0.06$, $\ln \frac{\hat{p}_0}{1-\hat{p}_0} = 0.03$, $\kappa = 0.019$.

learning is skewed towards the LLM's prior: for example, if $\hat{p}_0 > \frac{1}{2}$, we have $\bar{p} - \frac{1}{2} > \frac{1}{2} - \underline{p}$.

For the remainder of the equilibrium, the investor's value in the continuation region is given in (22). The two value matching conditions (24) and (25), and the two smooth pasting conditions (26) and (27) determine the constants C_1 , C_2 and the optimal policies $\underline{p}$, $\bar{p}$.

Figure 3 provides an illustration of the equilibrium. As shown in the left panel, the investor's attains the highest possible payoff, $U = 0$, when she is certain about her preference—that is, when $p = 0$ or 1 . If the investor perfectly learns ω , the LLM's posterior belief is also $\hat{p} = 0$ or 1 (an infinite z implies an infinite $\hat{z}$). However, the investor stops learning before reaching $p = 0$ or 1 in equilibrium. As illustrated by Figure 3, the investor has a higher incentive to communicate and stops at more extreme $\underline{p}, \bar{p}$ when the LLM partially learns from past communication.

3.3 Comparative Statics of AI Advising

We present a few comparative statics based on the baseline equilibrium in Proposition 2.

Communication cost c . The investor trades off the gain from continuing communication to learn about ω against the communication cost. As illustrated in Figure 4, when the communication cost c is larger, the principal learns less as suggested by a higher $\underline{p}$ and a lower $\bar{p}$ —she stops communication when she is less sure about ω .

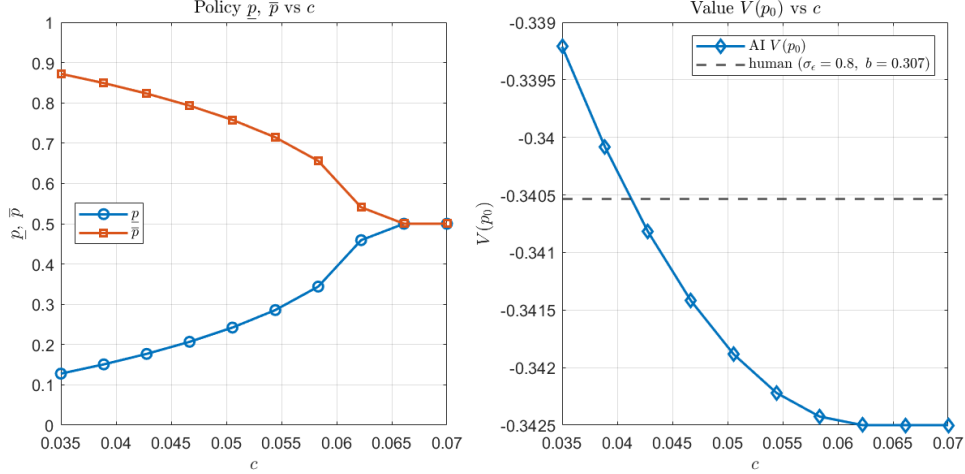


Figure 4: **The effects of communication cost c .** The left panel plots the optimal stopping thresholds $p, \bar{p}$ as a function of c , and the right panel shows the investor’s value at prior $V(p_0)$ as a function of c . Baseline parameters: $\mu_1 = 0.3$, $\mu_0 = 0$, $\sigma_\epsilon = 0.8$, $\sigma = 0.3$, $\lambda = 0.3$, $p_0 = \hat{p}_0 = 0.5$, $\kappa = 0$.

Preference uncertainty $\Delta\mu$. We discuss how preference uncertainty affects AI advising. The inefficient communication of soft information ω is the core friction when the investor consults an LLM.

First, we discuss how the magnitude of preference uncertainty affects the equilibrium. When the distributions of $\tilde{\theta}_1$ and $\tilde{\theta}_0$ are close to each other, preference uncertainty is less consequential, and receiving an impersonalized recommendation from the LLM causes little harm. Recall that $\tilde{\theta}_1 \sim N(\mu_1, \sigma_\epsilon^2)$ and $\tilde{\theta}_0 \sim N(\mu_0, \sigma_\epsilon^2)$ are independent normal random variables. In our numerical exercise, we fix the value of μ_1 and vary μ_0 for $\mu_0 < \mu_1$. In Figure 5, the higher is μ_0 —or the smaller is $\Delta\mu \equiv \mu_1 - \mu_0$, the investor learns less about ω (the optimal stopping thresholds become less extreme) and enjoys a higher value given the less consequential preference uncertainty. Intuitively, when the task is simple and personalization matters less, the investor has less incentive to figure out her exact preferences and she will still be better off.

Therefore, tasks without distinct contingencies are more suitable for AI advising. For example, tourism planning is easier, as the destination choice often reflects stereotypical preferences. In contrast, insurance or medical consultations involve much more soft information, where identifying the correct ω is critical.

We also examine the effects of the LLM’s pretraining model or its belief $p_t = \hat{p}_0$. As illustrated in Figure 6, the principal learns more information if the LLM’s belief is more extreme. Intuitively, if the LLM is trained towards a specific customer type, the investor benefits from learning more—either because she matches that type or needs to be more informed herself to adjust the action on her own.

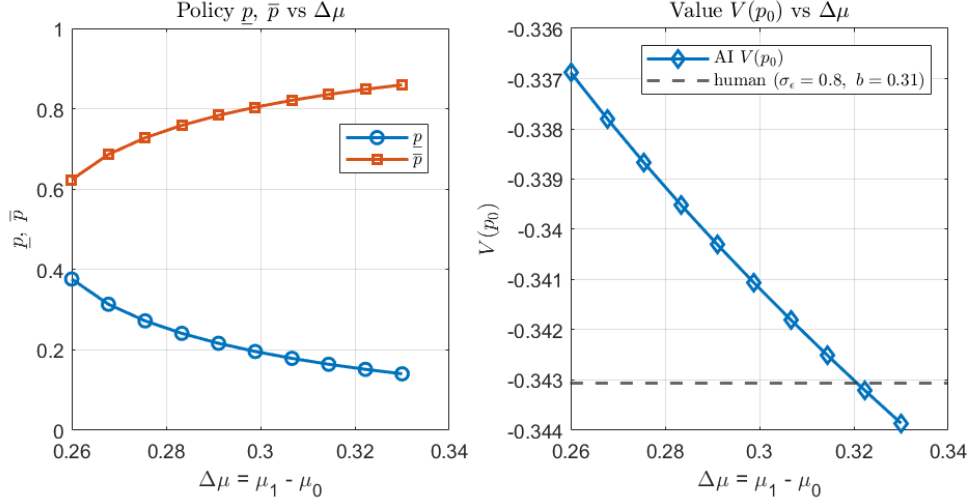


Figure 5: **The effects of preference uncertainty $\Delta\mu$.** The left panel plots the optimal stopping thresholds $\underline{p}, \bar{p}$ as a function of $\Delta\mu \equiv \mu_1 - \mu_0$, and the right panel shows the investor’s value at prior $V(p_0)$ as a function of $\Delta\mu$. Baseline parameters: $c = 0.045$, $\mu_0 = 0$, $\sigma_\epsilon = 0.8$, $\sigma = 0.3$, $\lambda = 0.3$, $p_0 = \hat{p}_0 = 0.5$, $\kappa = 0$.

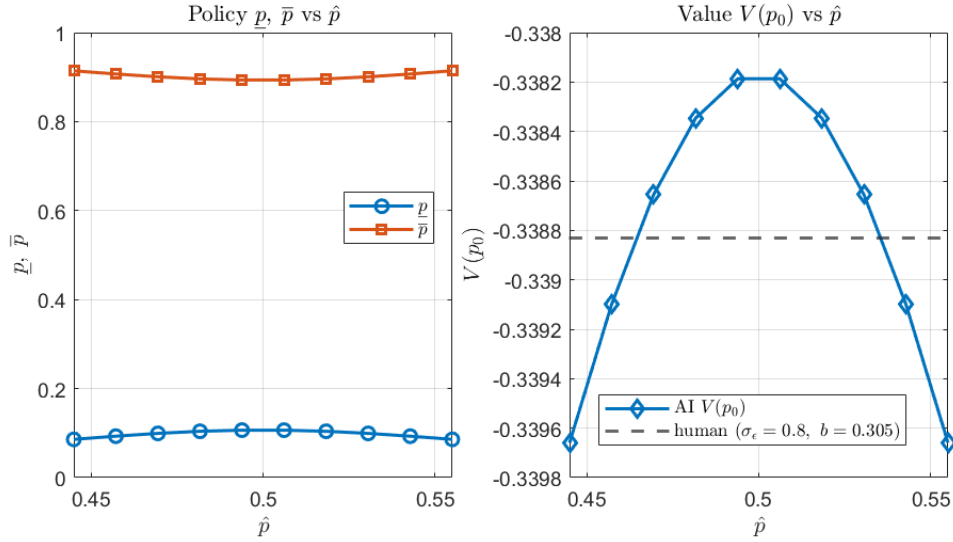


Figure 6: **The effects of the LLM’s belief $\hat{p}$.** The left panel plots the optimal stopping thresholds $\underline{p}, \bar{p}$ as a function of $\hat{p}$, and the right panel shows the investor’s value at prior $V(p_0)$ as a function of $\hat{p}$. Baseline parameters: $\mu_1 = 0.3$, $\mu_0 = 0$, $c = 0.045$, $\sigma_\epsilon = 0.8$, $\sigma = 0.3$, $\lambda = 0.3$, $p_0 = 0.5$, $\kappa = 0$.

Fundamental uncertainty σ_ϵ^2 . Last, we analyze the effects fundamental uncertainty, which is captured by the variance σ_ϵ^2 of the fundamental state, $\tilde{\theta}_1$ or $\tilde{\theta}_0$. When the advisor is a human, the investor’s utility is determined by the residual fundamental uncertainty—so she enjoys a higher payoff when σ_ϵ^2 is small.

When it comes to AI advising, the investor is still subject to residual fundamental uncertainty that arises from inefficient communication of preference ω , under which the LLM’s

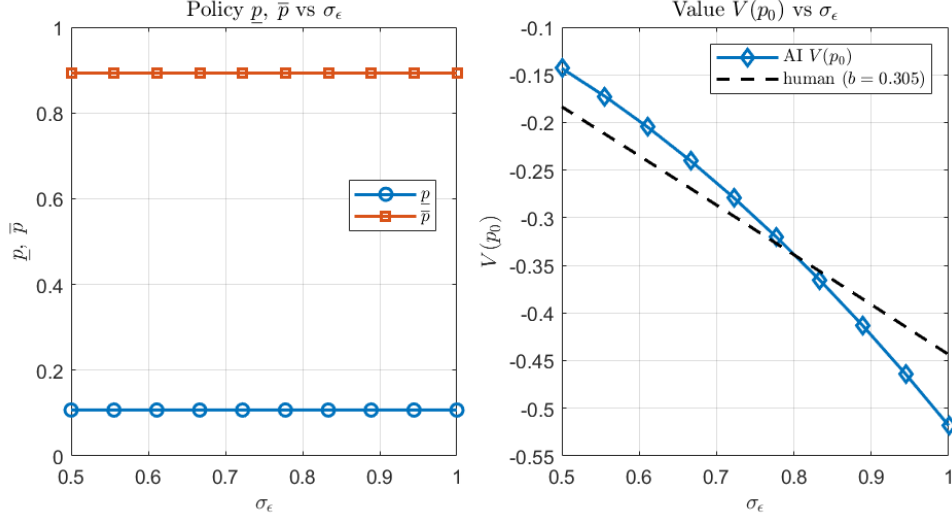


Figure 7: **The effects of fundamental uncertainty σ_ϵ .** The left panel plots the optimal stopping thresholds $\underline{p}, \bar{p}$ as a function of σ_ϵ , and the right panel shows the investor’s value at prior $V(p_0)$ as a function of σ_ϵ . Baseline parameters: $\mu_1 = 0.3$, $\mu_0 = 0$, $c = 0.045$, $\sigma = 0.3$, $\lambda = 0.3$, $p_0 = \hat{p} = 0.5$, $\kappa = 0$.

recommendation is a noisy signal of the true fundamental. Hence, in AI advising, the investor also enjoys a higher payoff when σ_ϵ^2 is small, as shown in the lower panel of Figure 7. Interestingly, as shown in the upper panel, the investor’s optimal policies, $\underline{p}$ and $\bar{p}$, do not vary with fundamental uncertainty. This suggests that the inefficiency in AI advising comes from preference uncertainty, and fundamental uncertainty does not affect the endogenous communication about preference ω .

4 Empirical Analysis

In this section, we discuss model implications and test them using prompt-based simulations in LLMs.

4.1 Testable Hypotheses

Building on the theoretical model and its comparative statics, we derive testable hypotheses that connect the model’s primitives to observable investor behavior and advice outcomes. Each hypothesis articulates a directional prediction, identifies the underlying mechanism from the theory, and outlines an empirical strategy for validation.

In a controlled setting, we can simulate the advising process with an LLM using prompt response logs and resulting portfolio recommendations to directly test several model predictions. These experiments allow us to modify conversation parameters and monitor outcomes

while keeping other factors constant. There are additional testable hypotheses motivated by the theory that require observational data (such as records of advisor choice and performance from a financial platform); we list those in Appendix B.1.

H1 (Investor Learning and Decision Quality): The primary value of interacting with a memory-less LLM advisor stems from the investor clarifying their own initially uncertain preferences. A deeper interaction allows the investor to reduce their own “preference uncertainty,” leading to a final investment decision that is better aligned with their true objectives, even if the LLM’s output remains generic.

The model posits that the investor is initially “confused” and does not fully understand her own needs summarized by ω . The dialogue with the LLM provides signals that allow the investor to update her own beliefs about her type, moving her closer to certainty—that is, her belief about ω , p approaches 1 or 0. Because we control the LLM’s information set to be the last prompt (short memory), its recommendation does not adapt to a deeper conversation. However, the now better-informed investor can make a more appropriate final decision based on her sharpened posterior belief.

H2 (Investor Impatience and Early Termination): Investors who face a higher opportunity cost of time break off LLM conversations sooner and accept portfolios less tailored to their stated preferences.

In our model, an impatient investor may experience a Poisson “impatience shock” that prompts them to end the communication early. This premature truncation of the dialogue results in the investor gathering less information, which in turn leads to a portfolio decision that is less tailored to the investor’s true preferences. To test this hypothesis, we examine the relationship between conversation length and the investor’s final portfolio choice.

H3 (Memory Augmentation and Advisor Performance): Providing the LLM advisor with a form of persistent memory about past interactions will improve its advice quality.

The theoretical motivation suggests that human advisors naturally retain and recall earlier parts of conversations and the client’s background, unlike a standard LLM which may lose mid-conversation context due to lack of memory. Enhancing the LLM with tools such as automatically generated summaries of past chat history should help reduce this information loss, enabling the LLM to better understand and address the investor’s needs.

4.2 Hypothesis Testing with LLM Simulations

To complement our analytical framework, we use prompt-based LLM simulations, a methodological innovation that enables the testing of complex economic theories in realistic, inter-

active settings. Unlike traditional laboratory experiments or analytical models, these simulations can generate dynamic, multi-turn conversations that closely approximate real-world advisory relationships. This approach bridges the critical gap between theoretical rigor and empirical realism, allowing us to observe how theoretical mechanisms like information acquisition and belief updating operate in practice. This section describes how we construct investor profiles, implement an LLM-based advisor, and map the discrete simulation environment back to the continuous-time model.

4.2.1 Data and Investor Profiles

The baseline for “optimal” advice in our simulations comes from the Vanguard Investor Questionnaire, a widely used 11-question survey that generates personalized asset-allocation recommendations based on an individual’s investment horizon, financial stability and risk tolerance.¹⁰ For example, the questionnaire recommends investors with short time horizons to hold a smaller fraction of equity in their portfolio while those with longer horizons to take on more risks. Similarly, it notes that a stable income stream allows investors to tolerate greater market volatility and therefore merits a higher equity weight.

To build our rule-based frictionless benchmark, we first simulate random investor profiles. For each of $n = 500$ hypothetical investors, we randomly select an answer for every question on the questionnaire, respecting the number of choices available for each item. The randomization uses a fixed seed to ensure replicability and produces a unique set of responses for each simulated investor.

We then take each simulated profile and use it to complete the Vanguard questionnaire on the website, recording the recommended stock/bond allocation for each hypothetical investor. Each profile in our benchmark is thus defined by: (i) a list of question-answer pairs that can be expressed in natural language, which correspond to the investor’s underlying preference ω , and (ii) the recommended allocation between equity vs fixed income, returned by the Vanguard algorithm, which maps to the ideal recommendation $\omega\theta_1 + (1-\omega)\theta_0$ without any frictions. In this way, we obtain a consistent set of simulated investors, grounded in established industry practice.

4.2.2 Conversation with the LLM Advisor

For each investor profile we simulate a conversation between the investor and an LLM-based advisor. The LLM we use is the then-state-of-the-art OpenAI GPT-5, accessed via

¹⁰The service is available at the website: <https://investor.vanguard.com/tools-calculators/investor-questionnaire/questions>. The questionnaires are listed in Appendix B.2.

an API at temperature 0.75. Conversations are orchestrated through system and developer messages that clearly specify each agent’s role and information set. The investor’s system prompt states that she is preparing to set up an investment portfolio and will interact with a financial advisor. It also includes the full text of her questionnaire profile. As explained below, the investor chats with the LLM based on this questionnaire but makes portfolio decisions without it. The advisor’s system prompt specifies that it is a financial advisor whose objective is to determine the client’s optimal allocation between equities and bonds. It is instructed to remain in information-gathering mode and to ask diagnostic questions about financial situation, investment experience, time horizon and risk preferences until the conversation ends. Importantly, the system prompt explicitly forbids the advisor from providing any recommendations before being told to do so, ensuring that all intermediate messages consist only of questions.

The interaction unfolds in discrete rounds:

1. **Eliciting the prior.** Before any questions are asked, the investor reports her prior preferred equity allocation. She is prompted to return her intended equity percentage as JSON, and this value is logged. This step parallels the initial belief p_0 in our theoretical model.
2. **Question–answer loop.** A developer message instructs the advisor to “ask your client one question that will help you identify their optimal split between equities and bonds.” The advisor generates a question, which we append to the conversation. The investor’s developer prompt tells her to answer concisely. She answers truthfully *according to her profile questionnaires*, and the answer is appended to both the investor and advisor message lists. In subsequent rounds the developer instructs the advisor to “ask your client another question,” so that exactly one question is asked per round. The conversation thus alternates between one question from the advisor and one answer from the investor.
3. **Termination decision.** After each round, the simulation consults the pre-drawn termination schedule (described below) to determine whether the conversation should stop. If the investor terminates, the question–answer loop ends; otherwise the advisor is prompted to ask another question.
4. **Recommendations.** Upon termination the advisor is asked to provide equity-allocation recommendations in two ways. First, to simulate the memoryless case, the advisor receives only the most recent question and answer and is instructed to return the optimal equity allocation for the individual. Second, to simulate memory, the advisor receives

the entire chat history (including all simulated questions and answers) and is instructed to produce an optimal equity allocation. In both cases, the advisor’s output must be a JSON snippet specifying the recommended equity percentage; no narrative explanation is allowed. For the full-information advisor, the advisor is given the complete investor profile up front and asked once for the optimal equity allocation.

5. **Final decision.** Finally, the investor is told that “the advisor recommends an equity allocation of $x\%$ ” and is prompted to choose her final allocation. Like the prior, the final allocation is returned as JSON. This step captures the investor’s own action in the theoretical model after observing the advisor’s message.

Throughout the conversation, we log every prompt and response. The strict formatting (e.g., JSON outputs, one question per round, no unsolicited recommendations) reflects an attempt to control the LLM’s behavior and reduce hallucinations. The advisor’s inability to update its belief in the memoryless condition comes from receiving only a single question–answer pair as input; this effectively resets its context window at every termination event.

We implement three variants of the LLM advisor to isolate the effect of memory:

- **Advisor without memory.** Motivated by the architecture of current LLMs, we simulate an advisor that cannot recall previous signals. After each prompt, the LLM uses only the most recent question–answer pair as context to recommend an allocation. This design captures the short context windows of Transformer models, whose self-attention complexity scales quadratically in the input length and cannot reliably capture the entire query history. Wang and Sun (2025) show that even when longer contexts are available, retrieval accuracy can deteriorate rapidly due to interference from earlier inputs; the probability of recalling the most recent key–value pair declines log-linearly as similar distractors accumulate. Our “no-memory” LLM therefore approximates our baseline model that the advisor updates its prior based on only the last signal, $\hat{p}_t = \mathbb{E}[\omega = 1 | ds_{t-}]$. (In the model, the advisor’s belief remains fixed at a pretraining prior under the continuous time limit.)
- **Advisor with memory.** In this scenario, the LLM is fed the entire conversation history as context. It has access to previous answers and can refine its recommendation as it learns more about the investor. This scenario is our best approximation of an unbiased human advisor with a transcript of the conversation. This aligns with our model extension with $\kappa = 1$: both the investor and the advisor update their belief based on the full history of signals $\{s_t\}$ before termination; if they share the same prior, then $\hat{p}_t = p_t$.

- **Advisor with full information.** This counterfactual LLM receives the investor’s entire questionnaire profile at once. It directly observes ω and effectively faces no preference uncertainty. This case can also mimic a human advisor who perfectly elicits soft information instantaneously, with no bias ($b = 0$). The full-information advisor’s recommendation provides an upper bound on achievable accuracy. Note that the advisor’s recommendation may still fall short of the frictionless Vanguard benchmark because of fundamental uncertainty. While the model assumes that advisors observe the fundamental realizations θ_1 and θ_0 , in practice, the LLM only observes a noisy signal of these fundamentals.

4.2.3 Termination and Cost of Communication

We cap the interaction at a maximum of eleven rounds to mirror the finite length of the Vanguard questionnaire. However, an investor does not necessarily complete all eleven rounds, as the dialogue can end sooner based on one of two mechanisms: exogenous or endogenous termination. To model impatience or external factors that cut a conversation short, we introduce an exogenous termination probability of 0.10 per round. Before the simulation begins, we conduct a Bernoulli draw for each round to determine if an “impatience shock” occurs. If a shock is scheduled for a given round, the conversation stops automatically after that round’s question and answer are complete. If no exogenous shock occurs, the investor agent makes an active, endogenous decision to continue or stop. After answering the advisor’s question, the investor is prompted with a decision frame that explicitly asks them to weigh the costs and benefits of more interaction: *“Your time is valuable, so each round of communication with the advisor carries a cost. You should choose to continue interacting with the advisor if and only if your expected informational gain from an additional round of communication exceeds your subjective cost of communication. Would you like to continue or terminate the conversation...?”* This prompt directly simulates the optimal stopping trade-off.

This dual-termination mechanism serves as a discrete analogue of the continuous-time optimal stopping problem in our theoretical model. Each question-answer round in the simulation corresponds to a small increment of time dt . The exogenous termination probability models the Poisson shock intensity λ , while the cost of continuing the conversation mirrors the flow cost c . The endogenous decision to terminate corresponds to the investor choosing the optimal stopping time τ once her belief p_t makes further interaction suboptimal. The theoretical problem is expressed as:

$$(SP) \quad \sup_{\tau \geq 0} \mathbb{E}^\omega \left\{ \int_0^\tau e^{-\lambda t} [-c + \lambda g(p_t, \hat{p}_t(p_t))] dt + e^{-\lambda \tau} g(p_\tau, \hat{p}_\tau(p_\tau)) \right\},$$

where p_t is the investor’s belief about her type and $g(\cdot)$ is the payoff from acting.

4.3 Empirical Results and Hypothesis Testing

We now present the empirical results from our LLM simulations to test the three main hypotheses derived from the theoretical model. Our analysis uses data from 500 simulated investor profiles, each with 5 iterations, yielding 2,500 total observations. The results provide strong support for the model’s predictions about the role of investor learning, the impact of communication costs, and the benefits of memory augmentation in AI advising.

4.3.1 Testing H1: Investor Learning and Decision Quality

The first hypothesis posits that LLM advising help investor clarify their own preferences. Table 1 presents the regression results testing this hypothesis. Across measures, the data show notable gains in investment accuracy that arise from the advisory process itself.

Investors improved significantly even before receiving any recommendations. As shown in Columns (1) and (2), accuracy increased by 14.1 percentage points before any advice, reaching 15.7 percentage points after the advisor’s input. The small 1.6 point boost directly from recommendations suggests most learning comes from interaction rather than specific advice.

Interaction intensity is also vital. Results from Columns (3) to (8) show that each additional round between advisor and investor increased interim accuracy by 0.986 points and final accuracy by 1.030 points. This implies iterative exchanges help investors gradually refine their understanding of risk and preferences.

The actual words exchanged drive learning as well. Column (5) shows that every additional word spoken by either party increases accuracy by 0.017 points. Columns (6) and (7) reveal that both advisor words and investor words improve accuracy, with investor input having a slightly larger impact. Articulating and reflecting on one’s own beliefs proves more effective than simply receiving advice. The multivariate model in Column (8) adds nuance to these effects. Once both advisor and investor word counts are included, advisor words remain a significant predictor of accuracy while investor words become insignificant. This suggests that while both parties contribute to the learning process, the advisor’s role may be more complementary than substitutive, helping to structure the conversation and guide the investor’s self-reflection rather than simply providing direct recommendations.

The consistently strong effects of self-reflection and dialogue, compared to the minor role of recommendations, suggest the main benefit of LLM advising is in helping investors understand themselves. For AI advisory systems, this means the greatest value lies in designing

Table 1: Investor Learning and Decision Quality

This table reports regression results examining factors influencing the accuracy of investors' investment choices, measured at interim and final stages, as well as accuracy improvements. Columns (1)-(2) display improvements in investment choice accuracy at pre- and post-recommendation stages. Column (3) specifically illustrates investors' interim investment choice accuracy, while Columns (4)-(8) present investors' final investment choice accuracy. Accuracy here is defined as 100 minus the deviation from the optimal portfolio allocation, expressed in percentage points. Independent variables include the number of interaction rounds (# Rounds, Total # Rounds), total words exchanged, and separately, the number of words contributed by the advisor and the investor. Standard errors, clustered by investor profile, are shown in parentheses. Profile fixed effects are incorporated in Columns (3)-(8) to control for investor-specific characteristics. ***, **, and * denote the 1%, 5%, and 10% confidence level, respectively.

	Accuracy Improvement		Interim	Investor's Accuracy				
	Pre-Rec	Post-Rec		Final				
	(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)
# Rounds			0.986*** (0.100)					
Total # Rounds				1.030*** (0.150)				
Total # Words					0.017*** (0.002)			
Total # Advisor's Words						0.027*** (0.004)		0.021*** (0.007)
Total # Investor's Words							0.040*** (0.006)	0.011 (0.011)
Constant	14.1*** (0.589)	15.7*** (0.682)						
Profile FE	N	N	Y	Y	Y	Y	Y	Y
Clustered by Profile	Y	Y	Y	Y	Y	Y	Y	Y
Observations	2,500	2,500	9,559	2,500	2,500	2,500	2,500	2,500
Adjusted R^2	-	-	0.521	0.790	0.790	0.790	0.789	0.790

Note:

* $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$

effective interaction and questioning frameworks to help users consider their choices, rather than focusing only on recommendation algorithms.

4.3.2 Testing H2: Investor Impatience and Early Termination

The second hypothesis examines whether investor impatience, manifested through early conversation termination, undermines the advisor’s ability to provide accurate recommendations. Table 2 presents regression results analyzing the factors that influence the accuracy of the AI advisor’s investment recommendations.

Table 2: Investor Impatience and Early Termination

This table reports regression results examining factors influencing the accuracy of the AI advisor’s investment recommendations. All columns present the accuracy of the advisor’s recommendations, measured as 100 minus the deviation from the optimal portfolio allocation, expressed in percentage points. Columns (1)-(4) and (6) examine individual factors in isolation, while Column (5) presents a specification including both advisor and investor word counts, and Column (7) includes the total number of interaction rounds and the exogenous termination indicator. Independent variables include the total number of interaction rounds (Total # Rounds), total words exchanged, separately the number of words contributed by the advisor and the investor, and an indicator for whether the conversation was terminated exogenously rather than reaching a natural conclusion. Standard errors, clustered by investor profile, are shown in parentheses. Profile fixed effects are incorporated in Columns (3)-(8) to control for investor-specific characteristics. ***, **, and * denote the 1%, 5%, and 10% confidence level, respectively.

	Accuracy of Advisor’s Recommendation						
	(1)	(2)	(3)	(4)	(5)	(6)	(7)
Total # Rounds	1.020*** (0.155)						0.872*** (0.187)
Total # Words		0.018*** (0.003)					
Total # Advisor’s Words			0.027*** (0.004)		0.018*** (0.007)		
Total # Investor’s Words				0.042*** (0.006)	0.016 (0.011)		
$\mathbb{I}\{\text{Exogenous Termination}\}$						-2.620*** (0.471)	-0.962* (0.567)
Profile FE	Y	Y	Y	Y	Y	Y	Y
Clustered by Profile	Y	Y	Y	Y	Y	Y	Y
Observations	2,500	2,500	2,500	2,500	2,500	2,500	2,500
Adjusted R^2	0.469	0.471	0.470	0.469	0.471	0.460	0.470

Note:

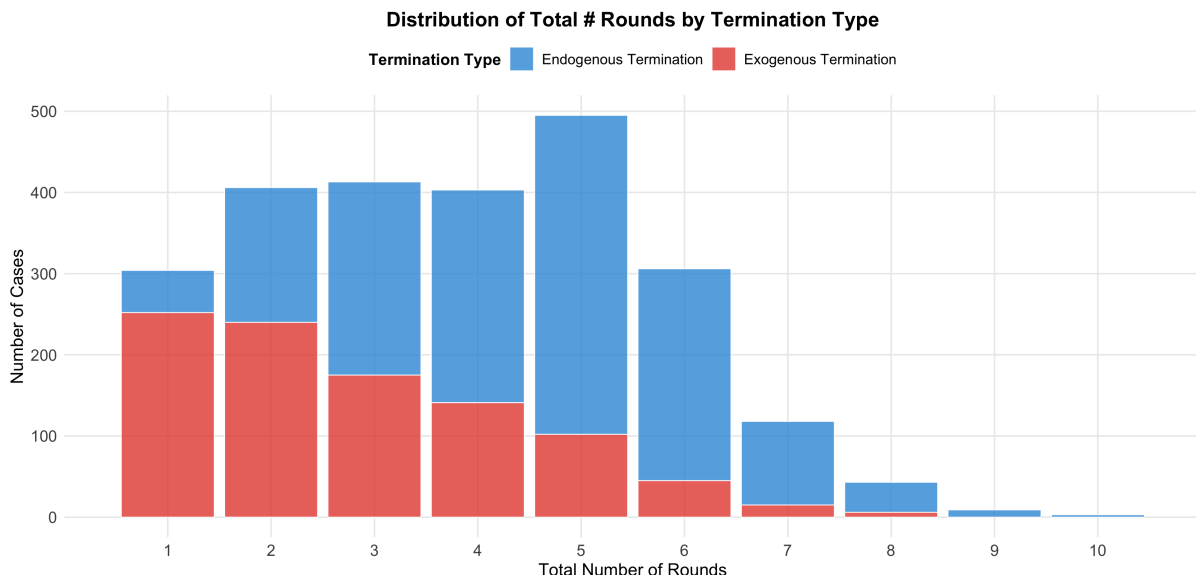
* $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$

The results demonstrate a strong positive relationship between interaction intensity and advisor recommendation accuracy. Column (1) shows that each additional round of interac-

tion increases the advisor’s recommendation accuracy by 1.020 percentage points, indicating that extended dialogue enables the advisor to better understand investor preferences and circumstances. This finding is reinforced by the word-count analyses in Columns (2)-(4).

The decomposition of word contributions reveals interesting dynamics in the advisory relationship. When examined separately, advisor words (Column 3) increase recommendation accuracy by 0.027 points per word, while investor words (Column 4) show an even larger effect of 0.042 points per word. However, when both are included simultaneously in Column (5), the advisor’s words remain highly significant while the investor’s words become statistically insignificant. This pattern, aligned with the findings in Table 1, suggests that while investor input is valuable, the advisor’s ability to process and respond to that information is the critical factor in generating accurate recommendations.

Figure 8: Endogenous and Exogenous Terminations



This stacked bar chart shows the distribution of conversation rounds by termination type across the dataset. The x-axis represents the total number of question rounds, while the y-axis shows the absolute count of cases. Each bar is divided into two categories: Endogenous Termination (blue) and Exogenous Termination (red). Endogenous termination occurs when the conversation naturally concludes based on the conversation flow, while exogenous termination happens due to external shocks.

The distribution of the total number of rounds in Figure 8 distinguishes between two types of conversation endings, highlighting key dynamics of information acquisition. Endogenous terminations occur when investors stop early because their beliefs have sufficiently converged, representing optimal stopping in response to the tradeoff between information gained and the costs incurred. In contrast, exogenous terminations happen either due to a

random stopping probability or upon reaching the eleven-question cap, independent of belief convergence, though the maximum length of 11 rounds is never actually reached. The distribution reveals that exogenous termination is more frequent in the first few rounds, while endogenous termination becomes the majority after the third round. This indicates that the investor actively weighs the tradeoff between gaining more information from conversation and the cost of communication. Moreover, just a few rounds of conversation are already quite helpful for the investor to become satisfied.

The central finding regarding investor impatience emerges from the termination analysis. Column (6) shows that exogenous termination, indicating that conversations ended artificially rather than reaching natural completion, reduces advisor recommendation accuracy by 2.620 percentage points. This substantial penalty demonstrates that premature conversation endings significantly impair the advisor’s performance. The effect remains significant even when controlling for interaction intensity in Column (7), where exogenous termination still reduces accuracy by 0.962 percentage points despite the inclusion of total rounds as a control variable.

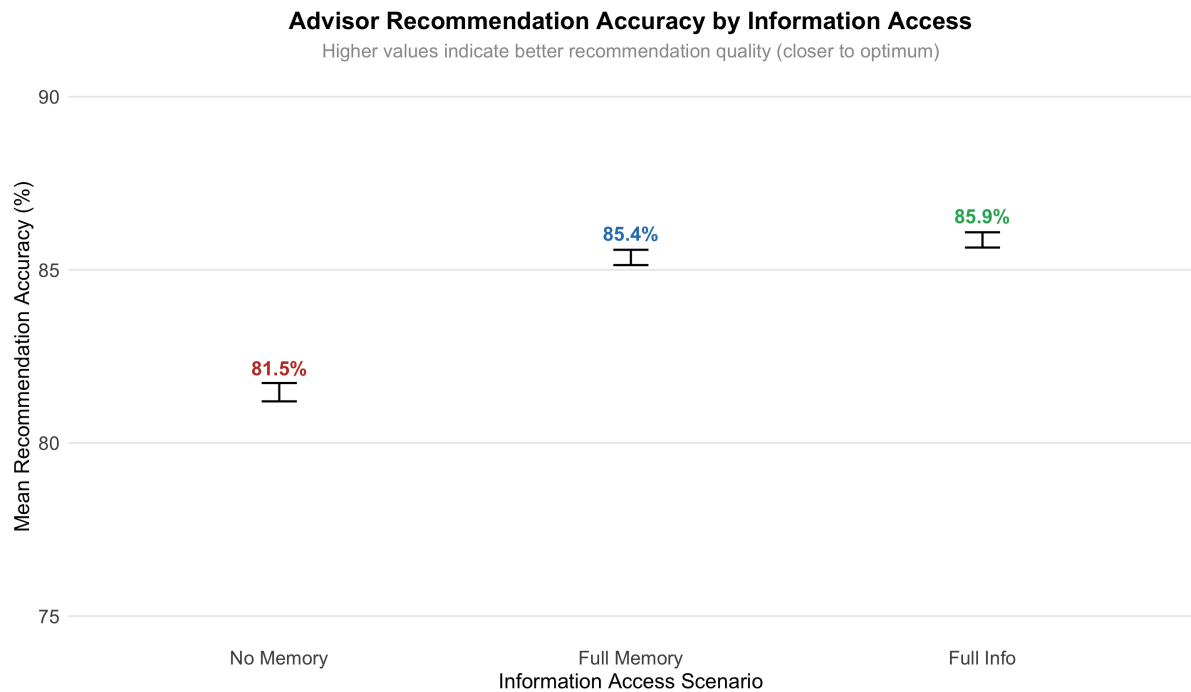
These results provide strong support for H2, indicating that investor impatience creates meaningful costs in advisory quality. The consistent negative impact of early termination, combined with the positive effects of extended interaction, suggests that the full advisory process requires adequate time and engagement to function effectively. For AI advisory systems, this highlights the importance of designing mechanisms that encourage sustained engagement and discourage premature exit from the advisory process.

4.3.3 Testing H3: Memory Augmentation and Advisor Performance

The third hypothesis examines whether providing LLM advisors with persistent memory improves their performance relative to memoryless systems. Figure 9 presents the results testing the benefits of memory augmentation.

The visualization demonstrates a clear hierarchy in recommendation quality based on information access: scenarios where the advisor has access to all available information including the investor’s complete profile achieve the highest accuracy, followed by scenarios where the advisor has access to the complete conversation history, while scenarios where the advisor only sees the current question perform the worst. This strongly supports Hypothesis 3, illustrating that both conversation memory and comprehensive profile information significantly enhance AI advisor performance. These results underscore that memory augmentation and full access to user information are both valuable for generating accurate investment recommendations.

Figure 9: Memory Augmentation and Advisor Performance



This bar chart compares AI advisor recommendation accuracy across three information access scenarios: No Memory (advisor only sees the current question), Full Memory (advisor has access to the complete conversation history), and Full Info (advisor has access to all available information including the investor’s complete profile). The y-axis shows mean recommendation accuracy as a percentage, with higher values indicating better recommendation quality and closer to the optimal equity allocation. Error bars represent standard errors.

4.3.4 Summary of Empirical Findings

The empirical results provide strong support for all three hypotheses derived from the theoretical model. H1 demonstrates that the primary value of LLM advising comes from investors clarifying their own preferences through self-reflection rather than receiving increasingly personalized recommendations, with each additional round improving accuracy by 1.03 percentage points and each word exchanged contributing 0.017 points to accuracy. H2 reveals that investor impatience, manifested through early conversation termination, significantly undermines advisor performance, with exogenous termination reducing recommendation accuracy by 2.62 percentage points. H3 shows that memory augmentation substantially improves AI advisor performance, with scenarios providing full information access achieving the highest accuracy, followed by full memory access, while no-memory scenarios perform worst.

The ability to generate large-scale, realistic simulations (2,500 conversations across 500 profiles) while maintaining the complexity of individual interactions provides unprecedented opportunities to test theories that depend on realistic communication patterns and adaptive behavior. By allowing researchers to observe mechanisms like optimal stopping in a practical setting, this methodology opens new possibilities for economic research in domains where human behavior is context-dependent and communication plays a central role. It offers a promising path forward for both theoretical development and the design of practical AI systems.

5 Conclusion

This study models the interaction between human and AI financial advisors as a dynamic optimal stopping game under two layers of uncertainty. We find that human advisors excel at interpreting soft, subjective information and clarifying ambiguous investor goals, while AI advisors provide unbiased, scalable recommendations but struggle to process unspoken preferences. This “soft information gap” limits the efficiency of AI-driven advice in complex decision-making contexts. Our results show that advisor value depends on the clarity of investor preferences. When goals are uncertain or evolving, human advisors’ interpretive strengths can outweigh their incentive biases. Conversely, when preferences are well-defined, unbiased AI advisors can match or outperform humans. Moreover, in contexts where human advice is heavily biased, AI guidance offers a clear advantage.

This work contributes a theoretical framework that integrates strategic communication, advisor incentives, and the role of soft information, extending classic models and complementing recent empirical findings on robo-advisors and generative AI. By formalizing the

inefficiencies created by digitizing soft information, we highlight key trade-offs in advisory relationships and set the stage for deeper integration of AI in finance.

A key methodological innovation of this study is the use of LLM simulations to operationalize and test our theoretical framework. By simulating realistic advisory interactions, we demonstrate how LLMs can serve as scalable, controlled environments for theory validation and refinement. This approach not only bridges theoretical and empirical analysis but also opens new avenues for using AI-driven simulations in behavioral finance and decision sciences.

Practically, our findings support a hybrid advisory approach. Novice or uncertain investors may benefit from initial human interaction, while experienced clients with clear goals can rely on AI platforms for efficient, unbiased advice. Enhancing AI systems with better memory and context-handling could further reduce information loss. Policymakers, meanwhile, should ensure that AI-driven advice remains transparent, unbiased, and aligned with fiduciary standards.

Future research could explore hybrid human–AI models, adaptive AI systems with extended memory, and applications in other fields such as medical or legal advising. Our framework provides a foundation for advancing the theory and practice of financial advice in the age of AI.

References

- Marianne Andries, Maxime Bonelli, and David Alexandre Sraer. Financial advisors and investors’ bias. *Available at SSRN 4691695*, 2024.
- Roland Bénabou and Guy Laroque. Using privileged information to manipulate markets: Insiders, gurus, and credibility. *The Quarterly Journal of Economics*, 107(3):921–958, 1992. doi: 10.2307/2118369.
- Jeremy Bertomeu and Iván Marinovic. A theory of hard and soft information. *The Accounting Review*, 91(1):1–20, 2016.
- Patrick Bolton, Xavier Freixas, and Joel Shapiro. Conflicts of Interest, Information Provision, and Competition in the Financial Services Industry. *Journal of Financial Economics*, 85(2):297–330, August 2007.
- Sean Cao, Wei Jiang, Junbo Wang, and Baozhong Yang. From man vs. machine to man+ machine: The art and ai of stock analyses. *Journal of Financial Economics*, 160:103910, 2024.
- Bruce Ian Carlin and Gustavo Manso. Obfuscation, Learning, and the Evolution of Investor Sophistication. *The Review of Financial Studies*, 24(3):754–785, March 2011.

- Ida Chak, Karen Croxson, Francesco D’Acunto, Jonathan Reuter, Alberto G Rossi, and Jonathan M Shaw. Improving household debt management with robo-advice. Technical report, National Bureau of Economic Research, 2022.
- John Chalmers and Jonathan Reuter. Is Conflicted Investment Advice Better Than No Advice? *Journal of Financial Economics*, 138(2):366–387, November 2020.
- Joanne Chen and Brandon Yueyang Han. Corporate ai backfire and potential remedies. *Available at SSRN*, 2024.
- Roberto Corrao. Mediation markets: The case of soft information. Technical report, Working Paper, 2023.
- Anna M Costello, Andrea K Down, and Mihir N Mehta. Machine+ man: A field experiment on the role of discretion in augmenting ai-based lending models. *Journal of Accounting and Economics*, 70(2-3):101360, 2020.
- Vincent P Crawford and Joel Sobel. Strategic information transmission. *Econometrica: Journal of the Econometric Society*, pages 1431–1451, 1982.
- Francesco D’Acunto, Nagpurnanand Prabhala, and Alberto G Rossi. The promises and pitfalls of robo-advising. *The Review of Financial Studies*, 32(5):1983–2020, 2019.
- Christian Fieberg, Lars Hornuf, David Streich, and Maximilian Meiler. Using Large Language Models for Financial Advice, August 2024.
- Nicola Gennaioli, Andrei Shleifer, and Robert Vishny. Money doctors. *The Journal of Finance*, 70(1):91–114, 2015.
- Fiona Greig, Tarun Ramadorai, Alberto G Rossi, Stephen P Utkus, and Ansgar Walther. Human financial advice in the age of automation. *Available at SSRN 4301514*, 2024.
- Yue Guo, Siliang Tong, Tao Chen, Subodha Kumar, and Shumiao Ouyang. The impact of generative artificial intelligence on individual manual investment decisions: Empirical evidence from mutual funds. *Nanyang Business School Research Paper*, (23-24), 2022.
- Zhiguo He, Jing Huang, and Cecilia Parlatore. Information span and credit market competition. Technical report, National Bureau of Economic Research, 2024.
- John J Horton. Large language models as simulated economic agents: What can we learn from homo silicus? Technical report, National Bureau of Economic Research, 2023.
- Enrique Ide and Eduard Talamas. Artificial intelligence in the knowledge economy. In *Proceedings of the 25th ACM Conference on Economics and Computation*, pages 834–836, 2024.
- Doron Levit and Nadya Malenko. Nonbinding voting for shareholder proposals. *Journal of Finance*, 66(5):1579–1614, 2011. doi: 10.1111/j.1540-6261.2011.01682.x.
- José María Liberti and Mitchell A Petersen. Information: Hard and soft. *Review of Corporate Finance Studies*, 8(1):1–41, 2019.
- Juhani Linnainmaa, Brian Melzer, Alessandro Previtero, and Stephen Foerster. Financial advisors and risk-taking. *Unpublished manuscript*, 2018.

- Nelson F Liu, Kevin Lin, John Hewitt, Ashwin Paranjape, Michele Bevilacqua, Fabio Petroni, and Percy Liang. Lost in the middle: How language models use long contexts. *arXiv preprint arXiv:2307.03172*, 2023.
- Fangzhou Lu, Lei Huang, and Sixuan Li. Chatgpt, generative ai, and investment advisory. *Available at SSRN 4519182*, 2023.
- Andrey Malenko and Nadya Malenko. Proxy advisory firms: The economics of selling information to voters. *The Journal of Finance*, 74(5):2441–2490, 2019. doi: 10.1111/jofi.12779.
- Hongwei Mo and Shumiao Ouyang. (Generative) AI in Financial Economics. *Available at SSRN 5287110*, 2025.
- Marco Ottaviani and Peter Norman Sørensen. Reputational cheap talk. *The RAND Journal of Economics*, 37(1):155–175, 2006. doi: 10.1111/j.1756-2171.2006.tb00010.x.
- Shumiao Ouyang, Hayong Yun, and Xingjian Zheng. Ai as decision-maker: Ethics and risk preferences of llms. *Available at SSRN 4851711*, 2024.
- Alberto G Rossi and Stephen Utkus. The diversification and welfare effects of robo-advising. *Journal of Financial Economics*, 157:103869, 2024.
- Jesper Rüdiger and Adrien Vigier. Learning about analysts. *Journal of Economic Theory*, 180: 304–335, 2019. doi: 10.1016/j.jet.2019.01.001.
- Jeremy C. Stein. Information production and capital allocation: Decentralized versus hierarchical firms. *The Journal of Finance*, 57(5):1891–1921, 2002. doi: <https://doi.org/10.1111/0022-1082.00483>. URL <https://onlinelibrary.wiley.com/doi/abs/10.1111/0022-1082.00483>.
- Chupei Wang and Jiaqiu Vince Sun. Unable to forget: Proactive Interference reveals working memory limits in llms beyond context length. *arXiv preprint arXiv:2506.08184*, 2025.

A Technical Appendices

A.1 Additional Details for Section 2.2

In this part, we provide omitted details for the communication with the LLM.

Derivation of Eq. (13) . Since $p_t(z_t) = \frac{e^{z_t}}{1+e^{z_t}}$, Ito's Lemma implies

$$\begin{aligned}
dp_t &= \left(\frac{e^{z_t}}{1+e^{z_t}} \right)'_{z_t} dz_t + \frac{1}{2} \left(\frac{e^{z_t}}{1+e^{z_t}} \right)''_{z_t} (dz_t)^2 \\
&= \frac{e^{z_t}}{(1+e^{z_t})^2} dz_t + \frac{1}{2} \frac{e^{z_t}}{(1+e^{z_t})^2} \left(1 - 2 \frac{e^{z_t}}{1+e^{z_t}} \right) (dz_t)^2 \\
&= p_t(1-p_t) \left[\frac{1}{\sigma^2} (p_t - \frac{1}{2}) dt + \frac{1}{\sigma} dB_t \right] + \frac{1}{2} p_t(1-p_t)(1-2p_t) \frac{dt}{\sigma^2} \\
&= \frac{p_t(1-p_t)}{\sigma^2} dB_t.
\end{aligned}$$

The third equation uses the evolvement of z_t in (12), $(dz_t)^2 = \frac{(dB_t)^2}{\sigma}$ and $p_t(z_t) = \frac{e^{z_t}}{1+e^{z_t}}$.

Derivation of $\hat{p}_t = \mathbb{E}(\omega = 1 | ds_{t-})$ (baseline). We use the following Binomial approximation of the process of Brownian motion. The investor observes the whole binomial tree whereas the LLM updates its belief only based on the last signal. Specifically, The probability that the particle goes up is

$$\pi = 0.5 + \frac{\omega}{2\sigma} \sqrt{\Delta t},$$

and the particle goes down is

$$1 - \pi = 0.5 - \frac{\omega}{2\sigma} \sqrt{\Delta t},$$

and the size of jump is

$$u = \sigma \sqrt{\Delta t}, \quad d = -\sigma \sqrt{\Delta t}.$$

For a particle that starts from zero, we have

$$\pi u + (1 - \pi) d = (2\pi - 1) u = \frac{\omega}{\sigma} \sqrt{\Delta t} \times \sigma \sqrt{\Delta t} = \omega \Delta t,$$

which means it has a drift of $\omega \Delta t$.

Suppose the particle goes up in the last round. The posterior belief of the LLM is

$$\hat{p}_t = \frac{\pi(\omega = 1) \hat{p}_0}{\pi(\omega = 1) \hat{p}_0 + \pi(\omega = 0)(1 - \hat{p}_0)} = \frac{\left(0.5 + \frac{1}{2\sigma} \sqrt{\Delta t}\right) \hat{p}_0}{\left(0.5 + \frac{1}{2\sigma} \sqrt{\Delta t}\right) \hat{p}_0 + 0.5(1 - \hat{p}_0)}.$$

Under the continuous time setting where $\Delta t \rightarrow 0$, the above equation becomes $\hat{p}_t = \lim_{\Delta t \rightarrow 0} \frac{\left(0.5 + \frac{1}{2\sigma} \sqrt{\Delta t}\right) \hat{p}_0}{\left(0.5 + \frac{1}{2\sigma} \sqrt{\Delta t}\right) \hat{p}_0 + 0.5(1 - \hat{p}_0)} = \hat{p}_0$.

A.2 Proof of Lemma 1

Proof. The investor understands that the recommendation m^L is based on the LLM's belief $\hat{p}$ and satisfies

$$m^L(\hat{p}) = \hat{p}\theta_1 + (1 - \hat{p})\theta_0. \quad (29)$$

The investor chooses her optimal action $a(m^L)$ to solve

$$\max_a \mathbb{E}[U(a, p, \tilde{\theta}_1, \tilde{\theta}_0)] = -p\mathbb{E}[(a - \tilde{\theta}_1)^2 | m^L] - (1 - p)\mathbb{E}[(a - \tilde{\theta}_0)^2 | m^L].$$

Her optimal action is

$$a(m^L) = p\mathbb{E}[\tilde{\theta}_1 | m^L] + (1 - p)\mathbb{E}[\tilde{\theta}_0 | m^L]. \quad (30)$$

From the investor's perspective, $\tilde{\theta}_1$, $\tilde{\theta}_0$ and m^L in (29) are normal random variables. Hence, the conditional expectation of states given recommendation m^L are

$$\mu_{\theta_1|m} \equiv \mathbb{E}[\tilde{\theta}_1 | m^L] = \mu_1 + \frac{\hat{p}}{\hat{p}^2 + (1 - \hat{p})^2} [m^L - \hat{p}\mu_1 - (1 - \hat{p})\mu_0], \quad (31)$$

$$\mu_{\theta_0|m} \equiv \mathbb{E}[\tilde{\theta}_0 | m^L] = \mu_0 + \frac{1 - \hat{p}}{\hat{p}^2 + (1 - \hat{p})^2} [m^L - \hat{p}\mu_1 - (1 - \hat{p})\mu_0]. \quad (32)$$

The investor's expected utility given any recommendation m^L is then

$$\begin{aligned} \mathbb{E}[U(a(m^L), p, \tilde{\theta}_1, \tilde{\theta}_0) | m^L] &= -a(m^L)^2 + 2a(m^L) [p\mu_{\theta_1|m} + (1 - p)\mu_{\theta_0|m}] - p\mathbb{E}[\tilde{\theta}_1^2 | m^L] - (1 - p)\mathbb{E}[\tilde{\theta}_0^2 | m^L] \\ &= - \underbrace{\left[p \text{Var}(\tilde{\theta}_1 | m^L) + (1 - p) \text{Var}(\tilde{\theta}_0 | m^L) \right]}_{MSE} - p(1 - p)(\mu_{\theta_1|m} - \mu_{\theta_0|m})^2. \end{aligned}$$

The conditional expectations $\mu_{\theta_1|m}$, $\mu_{\theta_0|m}$ are given in (31) and (32). The conditional variance $\text{Var}(\tilde{\theta}_1 | m^L) = \text{Var}(\tilde{\theta}_1) - \frac{\text{Cov}(\tilde{\theta}_1, m^L)^2}{\text{Var}(m^L)} = \frac{(1 - \hat{p})^2 \sigma_\epsilon^2}{\hat{p}^2 + (1 - \hat{p})^2}$ and $\text{Var}(\tilde{\theta}_0 | m^L) = \frac{\hat{p}^2 \sigma_\epsilon^2}{\hat{p}^2 + (1 - \hat{p})^2}$. Now we take expectation over $m^L(\hat{p})$ in (29) to calculate the investor's expected utility when she seeks recommendation,

$$\begin{aligned} g(p, \hat{p}) &= \mathbb{E}[\mathbb{E}[U(a(m^L), p, \tilde{\theta}_1, \tilde{\theta}_0) | m^L]] \\ &= -\sigma_\epsilon^2 \frac{p(1 - \hat{p})^2 + (1 - p)\hat{p}^2}{\hat{p}^2 + (1 - \hat{p})^2} - p(1 - p)\mathbb{E} \left\{ \mu_1 - \mu_0 + \frac{2\hat{p} - 1}{\hat{p}^2 + (1 - \hat{p})^2} [m^L - \hat{p}\mu_1 - (1 - \hat{p})\mu_0] \right\} \\ &= -\sigma_\epsilon^2 \frac{p(1 - \hat{p})^2 + (1 - p)\hat{p}^2}{\hat{p}^2 + (1 - \hat{p})^2} - p(1 - p) \left[(\mu_1 - \mu_0)^2 + \frac{(2\hat{p} - 1)^2 \sigma_\epsilon^2}{\hat{p}^2 + (1 - \hat{p})^2} \right]. \end{aligned}$$

Note that $g(p, \hat{p})$ is a quadratic function in p . □

A.3 Proof of Proposition 2

Proof. Step 1. We show that the principal continues if $p \in (\underline{p}, \bar{p})$ and stops to seek final recommendation if $p = \underline{p}$ or $\bar{p}$. In addition, $\underline{p} + \bar{p} = 1$ so that $\underline{p}$ and $\bar{p}$ are symmetric around $\frac{1}{2}$.

To see this, we consider the following adjusted value function

$$v(p) = V(p) - g(p), \quad (33)$$

which captures the option value of waiting. When waiting is strictly optimal, $v(p) > 0$. We rewrite the HJB in the continuation region, (22), and boundary conditions in terms of v . Plugging $V(p) = v(p) + g(p)$ in (22) we have

$$\lambda v(p) = \frac{p^2(1-p)^2}{2\sigma^2} g''(p) - c + \frac{p^2(1-p)^2}{2\sigma^2} v''(p).$$

Importantly, $g(p)$ is a quadratic function of p so $g''(p)$ is a constant.

First, note that both the flow benefit of waiting $\frac{p^2(1-p)^2}{2\sigma^2} g''(p) - c$ and the volatility $\frac{p^2(1-p)^2}{2\sigma^2}$ are larger when p is close to $\frac{1}{2}$ and smaller when p is close to 0 or 1. Hence, we conjecture that the principal continues if $p \in (\underline{p}, \bar{p})$ and stops to seek final recommendation if $p = \underline{p}$ or $\bar{p}$. The boundary conditions in terms of $v(\cdot)$ are

$$v(\underline{p}) = 0, v'(\underline{p}) = 0, v(\bar{p}) = 0, v'(\bar{p}) = 0.$$

Second, both the flow benefit of waiting $\frac{p^2(1-p)^2}{2\sigma^2} g''(p) - c$ and the volatility $\frac{p^2(1-p)^2}{2\sigma^2}$ are symmetric for p around 0.5. Intuitively, the underlying $\theta = 0$ and 1 indicates borrower type (matching $\tilde{\theta}_i$) and are interchangeable. Therefore, $\underline{p} + \bar{p} = 1$. Also, the symmetry means that there exists a critical upper bound learning cost $c < \bar{c}$ under which $v(\frac{1}{2}) > 0$ and the above waiting solution holds.

Step 2. We solve for the differential equation in closed form. We rewrite the differential equation in log-likelihood ratio $z = \ln \frac{p}{1-p}$, so that $p(z) = \frac{e^z}{1+e^z}$, $1 - p(z) = \frac{1}{1+e^z}$. Note that

$$p'(z) = \frac{1}{z'(p)} = \frac{1}{\frac{1}{p} + \frac{1}{1-p}} = p(1-p). \quad (34)$$

Define $W(z) \equiv V(p(z))$. Then

$$W_z = V'(p)p'(z) = V'(p)p(1-p), \quad (35)$$

or equivalently $V'(p) = \frac{W_z}{p(1-p)}$. Then second derivative

$$V''(p) = \frac{d}{dp} \left[\frac{W_z}{p(1-p)} \right] = z'(p) \frac{d}{dz} \left[\frac{W_z}{p(1-p)} \right] = \frac{1}{p^2(1-p)^2} [W_{zz} - (1-2p)W_z].$$

Hence, the differential equation (22) becomes

$$-c + \lambda(g(p) - W(z)) + \frac{1}{2\sigma^2} [W_{zz} - (1-2p)W_z] = 0. \quad (36)$$

Now we substitute $W(z) = u(z)\hat{W}(z)$ and choose $u(z)$ to remove the W_z term in (36). Then

$$W_z = u_z \hat{W} + u \hat{W}_z, \quad W_{zz} = u_{zz} \hat{W} + 2u_z \hat{W}_z + u \hat{W}_{zz},$$

Using them in (36), we have:

$$-c + \lambda[g(p(z)) - u\hat{W}] + \frac{1}{2\sigma^2} \left\{ u\hat{W}_{zz} + [2u_z - (1-2p)u] \hat{W}_z + [u_{zz} - (1-2p(z))u_z] \hat{W} \right\} = 0.$$

To eliminate the $\hat{W}_z$ term, impose

$$2u_z - (1-2p)u = 0 \Leftrightarrow \frac{u_p}{u} = \frac{1-2p}{p'(z)} = \frac{1}{2} \left(\frac{1}{p} - \frac{1}{1-p} \right). \quad (37)$$

Integrate w.r.t. p , we need to set

$$u(z) = \sqrt{p(z)(1-p(z))},$$

and then the differential equation becomes

$$\hat{W}_{zz} + \left[\frac{u_{zz} - (1-2p)u_z}{u} - 2\lambda\sigma^2 \right] \hat{W} + 2\sigma^2[\lambda g(p) - c] = 0. \quad (38)$$

From the definition of u function in (37), we know that $\frac{u_z}{u} = \frac{1-2p}{2}$. Then

$$\frac{d}{dz} \left(\frac{u_z}{u} \right) = \frac{u_{zz}}{u} - \left(\frac{u_z}{u} \right)^2,$$

which implies the coefficient in front of $\hat{W}$ in (38) is a constant:

$$\begin{aligned}\frac{u_{zz}}{u} - \frac{(1-2p)u_z}{u} &= p'(z) \frac{d}{dp} \left(\frac{u_z}{u} \right) + \left(\frac{u_z}{u} \right)^2 - \frac{(1-2p)u_z}{u} \\ &= p(1-p) \cdot (-1) + \left(\frac{1-2p}{2} \right)^2 - \frac{(1-2p)^2}{2} = -\frac{1}{4}.\end{aligned}$$

Therefore, with the substitution of $z \equiv \ln \frac{p}{1-p}$ and $\hat{W}(z) \equiv \frac{W(z)}{\sqrt{p(z)(1-p(z))}} = \frac{V(p(z))}{\sqrt{p(z)(1-p(z))}}$, the original differential equation in (22) could be rewritten as

$$\hat{W}_{zz} - \left(\frac{1}{4} + 2\lambda\sigma^2 \right) \hat{W} + 2\sigma^2 [\lambda g(p) - c] = 0.$$

Let $\gamma = \sqrt{\frac{1}{4} + 2\sigma^2\lambda}$. The homogeneous solution of the above differential equation is

$$\hat{W}(z) = Ae^{z\gamma} + Be^{-z\gamma}.$$

Hence, using $e^z = \frac{p}{1-p}$, the homogeneous solution of the original differential equation (22) is

$$V(p) = W(p(z)) = u(p)\hat{W}(z) = C_1 p^{\gamma+\frac{1}{2}}(1-p)^{-\gamma+\frac{1}{2}} + C_2 p^{-\lambda+\frac{1}{2}}(1-p)^{\lambda+\frac{1}{2}}.$$

As for the particular solution,

$$Q_0(p) = q_0 - \frac{c}{\lambda} + (q_1 + q_2)p + \frac{\sigma^2 \lambda q_2}{2\gamma} \left[p^{\frac{1}{2}-\gamma}(1-p)^{\frac{1}{2}+\gamma} B_p \left(\gamma + \frac{1}{2}, \frac{1}{2} - \gamma \right) + p^{\frac{1}{2}+\gamma}(1-p)^{\frac{1}{2}-\gamma} B_{1-p} \left(\gamma + \frac{1}{2}, \frac{1}{2} - \gamma \right) \right] \quad (39)$$

where $B_z(a, b) = \int_0^z t^{a-1}(1-t)^{b-1}dt$ is the beta integral, and q_0, q_1, q_2 are the coefficients for stopping value $g(p) = q_0 + q_1 p + q_2 p^2$ in (19).

Step 3. Last, we show that $A = B$ from the boundary conditions and the symmetric property that $\underline{p} + \bar{p} = 1$.

□

B Appendices for Empirical Exercise

B.1 Hypotheses Testable via Observational Data

The second set of hypotheses involves predictions that can be examined using real-world observational data, such as records from a financial platform offering both human and LLM-based advising. In our context, Yingmi Wealth's Qieman (meaning "Hold On"), which is an intelligent investment advisory platform in China, provides a useful example: some investors

on the platform may choose to consult a human financial advisor, while others use an AI advisor. By leveraging data on these choices, alongside investor characteristics and outcomes, we can potentially test how the model’s mechanisms play out in practice.

H4 (Preference Uncertainty and Advisor Choice): Investors who are more uncertain about their own financial goals or risk preferences are more likely to choose a human advisor over an LLM-based advisor.

The theoretical setup suggests that when investors are unsure about their preferences, such as unclear risk tolerance or retirement goals, human advisors may be better equipped to interpret nuanced soft information through interactive conversation compared to AI-based tools. This advantage is further supported by the model’s ability to account for uncertainty in understanding investor preferences, which suggests that human advisors are more effective in dynamically clarifying the needs of investors with less self-awareness.

To test this hypothesis, we propose analyzing observational data on advisor selection based on investors’ self-reported uncertainty. One approach is to use survey or onboarding data to construct a metric of ex ante preference uncertainty. This metric could be derived from variability in responses to risk tolerance questions or the absence of a clear investment goal. Using such data, a regression or discrete-choice model could be estimated to determine the probability of selecting a human advisor versus an AI advisor as a function of this uncertainty metric. The prediction is that investors with greater preference uncertainty will exhibit a significantly higher likelihood of opting for human advisors.

H5 (Extreme Investor Types and AI Adoption): Investors with extreme initial risk profiles are more likely to opt for the LLM advisor, especially those whose risk assessment falls in the very conservative or very aggressive ends of the distribution.

The intuition from the model suggests that for extreme investor types, those who are either highly risk-averse or highly risk-seeking, the potential downside from any misalignment in the AI’s recommendations is relatively smaller. This arises because even if the AI’s recommendations are slightly off, they will still be close to the investor’s true preferences or optimal strategy. On the other hand, human advisors with biases—such as those stemming from sales commissions—introduce additional "cheap-talk" costs. This inherent impartiality of the LLM makes its advice more appealing to investors situated at the spectrum’s extremes.

To test this hypothesis, empirical analysis can be conducted by studying how investors’ risk profiles correlate with their choice of advisor. Specifically, using data on initial risk scores, such as scores derived from risk assessment questionnaires, we can examine whether investors who fall into the lowest or highest quantiles of risk tolerance are more likely to opt for the LLM over human advisors. A probit or logit regression can be used where the decision to adopt the LLM is regressed on indicators for "very low risk tolerance" and "very high risk

tolerance," controlling for other relevant factors. Positive coefficients on these indicators would support the hypothesis.

H6 (Commission Incentives and Advisor Choice): The adoption of the LLM advisor will be higher in settings where human financial advisors are paid on commission, as opposed to a fee or salary basis.

The model suggests that when human advisors have incentives to influence clients, such as earning commissions from selling specific products, their advice becomes more biased and less credible. In such situations, the informational value of human advisors' cheap-talk messages decreases, leading investors to favor the unbiased AI advisors. To test this theory, researchers could explore variations in compensation structures either across different regions or over time. For example, comparing adoption rates of LLM advisors in settings where some branches of a financial platform rely mainly on commission-based pay while others use fixed salaries.

A regression analysis can be employed with the proportion of investors opting for AI advisors as the dependent variable. The primary independent variable would measure the strength of commission-based incentives for human advisors. It is expected that stronger commission incentives or a shift toward a commission-heavy compensation model would correlate with increased AI advisor adoption. Establishing that investors choose LLM advisors significantly more often in high-commission contexts would offer evidence supporting the hypothesis.

H7 (Investor Experience and Advisor Performance): Among more experienced or financially sophisticated investors, those who effectively know their type with greater precision, the performance gap between the LLM and a human advisor is smaller and may even reverse in favor of the LLM.

The idea is that experienced investors can convey their objectives and constraints more clearly or already understand them well, so a human advisor's ability to uncover soft information becomes less critical. According to the model, when the investor's type is already known or obvious, the human advisor's traditional edge in interpreting the investor's needs vanishes, leaving only the downside of the human's potential bias versus the AI's objectivity. In such cases, the neutral LLM advisor could perform just as well or better in aligning recommendations with the investor's true preferences.

To test this hypothesis, observational data can be analyzed, focusing on variables related to investor experience, such as years of investment experience, trading volume, or financial literacy scores. This data would also need measures of advice outcomes, like realized portfolio returns, risk-adjusted performance, or consistency with stated goals. An empirical strategy could involve a regression analysis, where an interaction term between investor experience

and advisor type measures the relative effectiveness of AI and human advisors. Specifically, the hypothesis predicts that the interaction term ($\text{Experience} \times \text{LLM}$) will have a positive coefficient in the regression, indicating that the performance of AI advisors improves as investor experience increases.

B.2 Vanguard Questionnaire

1. Once you start withdrawing money from your investments, you plan to spend it over a period of...
 - (a) 2 years or less
 - (b) 3-5 years
 - (c) 6-10 years
 - (d) 11-15 years
 - (e) More than 15 years
2. When making a long-term investment, you plan to keep the money invested for...
 - (a) 1-2 years
 - (b) 3-4 years
 - (c) 5-6 years
 - (d) 7-8 years
 - (e) More than 8 years
3. When it comes to investing in stock or bond mutual funds or ETFs (or individual stocks or bonds) you would describe yourself as...
 - (a) Very inexperienced
 - (b) Somewhat inexperienced
 - (c) Somewhat experienced
 - (d) Experienced
 - (e) Very experienced
4. You plan to begin taking money from your investments in...
 - (a) 1 year or less

- (b) 1-2 years
 - (c) 3-5 years
 - (d) 6-10 years
 - (e) 11-15 years
 - (f) More than 15 years
5. Your current and future income sources (for example, salary, social security, pensions) are...
- (a) Very unstable
 - (b) Unstable
 - (c) Somewhat stable
 - (d) Stable
 - (e) Very stable
6. From September 2008 through October 2008, bonds lost 4%. If you owned a bond investment that lost 4% in two months, you would...
- (a) Sell all the remaining investment
 - (b) Sell a portion of the remaining investment
 - (c) Hold onto the investment and sell nothing
 - (d) Buy more of the remaining investment
7. The below table shows the greatest 1-year loss and the highest 1-year gain on 3 different hypothetical investments of \$10,000.
- Investment A (gain \$593; loss -\$164)
 - Investment B (gain \$1,921; loss -\$1,020)
 - Investment C (gain \$4,229; loss -\$3,639)

Given the potential gain or loss in any 1 year, you would invest your money in...

- (a) minimal volatility
- (b) moderate volatility
- (c) most volatility

8. During market declines, you tend to sell portions of your riskier assets and invest the money in safer assets. **(R)**
- (a) Strongly disagree
 - (b) Disagree
 - (c) Somewhat agree
 - (d) Agree
 - (e) Strongly agree
9. You would invest in a mutual fund or ETF (exchange-traded fund) based solely on a brief conversation with a friend, co-worker, or relative. **(R)**
- (a) Strongly disagree
 - (b) Disagree
 - (c) Somewhat agree
 - (d) Agree
 - (e) Strongly agree
10. From September 2008 through November 2008, stocks lost over 31%. If you owned a stock investment that lost about 31% in three months, you would...
- (a) Sell all the remaining investment
 - (b) Sell a portion of the remaining investment
 - (c) Hold onto the investment and sell nothing
 - (d) Buy more of the remaining investment
11. Generally, you prefer an investment with little or no ups and downs in value, and you are willing to accept the lower returns these investments may make. **(R)**
- (a) Strongly disagree
 - (b) Disagree
 - (c) Somewhat agree
 - (d) Agree
 - (e) Strongly agree

Note: Questions marked with (R) are reverse-scored items where higher numerical responses indicate lower risk tolerance.