

Soft Information, Hard Decisions: AI Advising

Jing Huang

Texas A&M University

Shumiao Ouyang

Oxford University

NBER-SAIF AI and Financial Markets

The AI Adoption Paradox

- ▶ AI advisors: no conflicts of interest, scalable, low cost

The AI Adoption Paradox

- ▶ AI advisors: no conflicts of interest, scalable, low cost
- ▶ Yet robo-advisors hold $< 2\%$ of global AUM
- ▶ Why do AI advisors underperform?

The AI Adoption Paradox

- ▶ AI advisors: no conflicts of interest, scalable, low cost
- ▶ Yet robo-advisors hold $< 2\%$ of global AUM
- ▶ Why do AI advisors underperform?

Answer: Soft Information Frictions

AI doesn't know user's preference $\Rightarrow$ User must debias AI's recommendation

Communication of Soft Information

- ▶ Traditionally collected via human interactions.
- ▶ Now conversational exchanges with AIs are possible.
 - ▶ Hardening soft information (He Huang Parlatores 2025)

Communication of Soft Information

- ▶ Traditionally collected via human interactions.
- ▶ Now conversational exchanges with AIs are possible.
 - ▶ Hardening soft information (He Huang Parlatores 2025)
- ▶ Still, communication of soft info is **inefficient** with AI.
 1. Prompt quality: investors cannot articulate soft traits.
 2. AI memory: synthesis across the conversation.

This Paper

- ▶ **Model:** Preference uncertainty + cheap talk (CS 1982).
 - ▶ **Key friction:** AI recommends based on its belief of preference; user debiases
- ▶ Communication about preference: Brownian learning+optimal stopping
 - ▶ Prompt quality and AI memory.

This Paper

- ▶ **Model:** Preference uncertainty + cheap talk (CS 1982).
 - ▶ **Key friction:** AI recommends based on its belief of preference; user debiases
- ▶ Communication about preference: Brownian learning+optimal stopping
 - ▶ Prompt quality and AI memory.
- ▶ **Implications:**
 - ▶ Prompt quality can substitute for memory
 - ▶ Memory benefits depend on prior alignment
 - ▶ **Opinionated AI can be optimal**

This Paper

- ▶ **Model:** Preference uncertainty + cheap talk (CS 1982).
 - ▶ **Key friction:** AI recommends based on its belief of preference; user debiases
- ▶ Communication about preference: Brownian learning+optimal stopping
 - ▶ Prompt quality and AI memory.
- ▶ **Implications:**
 - ▶ Prompt quality can substitute for memory
 - ▶ Memory benefits depend on prior alignment
 - ▶ **Opinionated AI can be optimal**
- ▶ **Empirics:** LLM simulations with SCF investor profiles
 - ▶ Two-sided simulation: both advisor and investor are LLM-generated
 - ▶ **Control of information set.**

Model

Setting: Investor

- ▶ **Preference uncertainty** (novel)

- ▶ Investor's preference type $\omega \in \{0, 1\}$

- ▶ $\omega = 1$: target is $\tilde{\theta}_1$ (equity); $\omega = 0$: target is $\tilde{\theta}_0$ (fixed income)

Setting: Investor

- ▶ **Preference uncertainty (novel)**
 - ▶ Investor's preference type $\omega \in \{0, 1\}$
 - ▶ $\omega = 1$: target is $\tilde{\theta}_1$ (equity); $\omega = 0$: target is $\tilde{\theta}_0$ (fixed income)
- ▶ Beliefs: $p = \mathbb{P}^{investor}(\omega = 1)$, $\hat{p} = \mathbb{P}^{AI}(\omega = 1)$
 - ▶ ω is only learnt through communication with AI

Setting: Investor

- ▶ **Preference uncertainty (novel)**
 - ▶ Investor's preference type $\omega \in \{0, 1\}$
 - ▶ $\omega = 1$: target is $\tilde{\theta}_1$ (equity); $\omega = 0$: target is $\tilde{\theta}_0$ (fixed income)
- ▶ Beliefs: $p = \mathbb{P}^{investor}(\omega = 1)$, $\hat{p} = \mathbb{P}^{AI}(\omega = 1)$
 - ▶ ω is only learnt through communication with AI
- ▶ Uncertainty about fundamental realizations
 - ▶ $\tilde{\theta}_1 \sim N(\mu_1, 1)$, $\tilde{\theta}_0 \sim N(\mu_0, 1)$

Setting: Investor

- ▶ **Preference uncertainty (novel)**

- ▶ Investor's preference type $\omega \in \{0, 1\}$

- ▶ $\omega = 1$: target is $\tilde{\theta}_1$ (equity); $\omega = 0$: target is $\tilde{\theta}_0$ (fixed income)

- ▶ Beliefs: $p = \mathbb{P}^{investor}(\omega = 1)$, $\hat{p} = \mathbb{P}^{AI}(\omega = 1)$

- ▶ ω is only learnt through communication with AI

- ▶ Uncertainty about fundamental realizations

- ▶ $\tilde{\theta}_1 \sim N(\mu_1, 1)$, $\tilde{\theta}_0 \sim N(\mu_0, 1)$

- ▶ **Investor's utility:**

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

AI Recommendation

- ▶ AI observes (θ_1, θ_0) and recommends:

$$m = \hat{p}\theta_1 + (1 - \hat{p})\theta_0$$

AI Recommendation

- ▶ AI observes (θ_1, θ_0) and recommends:

$$m = \hat{p}\theta_1 + (1 - \hat{p})\theta_0$$

- ▶ **User debiases if $\hat{p} \neq p$**
 - ▶ Tilt action a towards p
 - ▶ Update beliefs about θ_ω from m

AI Recommendation

- ▶ AI observes (θ_1, θ_0) and recommends:

$$m = \hat{p}\theta_1 + (1 - \hat{p})\theta_0$$

- ▶ **User debiases if $\hat{p} \neq p$**
 - ▶ Tilt action a towards p
 - ▶ Update beliefs about θ_ω from m

Investor Value when Seeking Recommendation

$$g(p, \hat{p}) = \underbrace{-\frac{(p - \hat{p})^2}{\hat{p}^2 + (1 - \hat{p})^2}}_{\text{debiasing cost}} - \underbrace{p(1 - p)(\Delta_\mu^2 + 2)}_{\text{preference uncertainty}}$$

Investor Value when Seeking Recommendation

$$g(p, \hat{p}) = \underbrace{-\frac{(p - \hat{p})^2}{\hat{p}^2 + (1 - \hat{p})^2}}_{\text{debiasing cost}} - \underbrace{p(1 - p)(\Delta_\mu^2 + 2)}_{\text{preference uncertainty}}$$

► Debiasing cost:

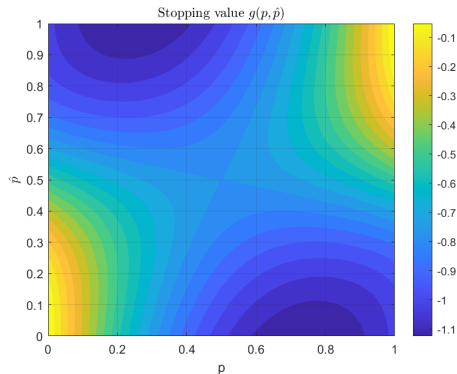
- Zero when $p = \hat{p}$
- Extreme $\hat{p} \Rightarrow$ more informative

Investor Value when Seeking Recommendation

$$g(p, \hat{p}) = \underbrace{-\frac{(p - \hat{p})^2}{\hat{p}^2 + (1 - \hat{p})^2}}_{\text{debiasing cost}} - \underbrace{p(1 - p)(\Delta_\mu^2 + 2)}_{\text{preference uncertainty}}$$

► Debiasing cost:

- Zero when $p = \hat{p}$
- Extreme $\hat{p} \Rightarrow$ more informative

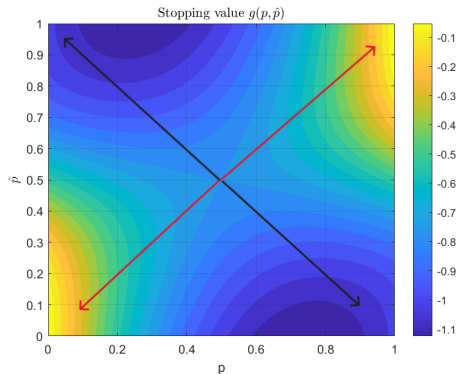


Investor Value when Seeking Recommendation

$$g(p, \hat{p}) = \underbrace{-\frac{(p - \hat{p})^2}{\hat{p}^2 + (1 - \hat{p})^2}}_{\text{debiasing cost}} - \underbrace{p(1 - p)(\Delta_\mu^2 + 2)}_{\text{preference uncertainty}}$$

► Debiasing cost:

- Zero when $p = \hat{p}$
- Extreme $\hat{p} \Rightarrow$ more informative



Communication about ω : Brownian Learning

- ▶ Beliefs $p, \hat{p}$ endogenous from communication. Public signal ds_t :

$$ds_t = \omega dt + \sigma dB_t$$

Prompt quality: $\frac{1}{\sigma}$

- ▶ Investor determines when to stop and seek m .
 - ▶ Cost $c dt$, impatience shock λ .

Communication about ω : Brownian Learning

- ▶ Beliefs $p, \hat{p}$ endogenous from communication. Public signal ds_t :

$$ds_t = \omega dt + \sigma dB_t$$

Prompt quality: $\frac{1}{\sigma}$

- ▶ Investor determines when to stop and seek m .
 - ▶ Cost cdt , impatience shock λ .

- ▶ **Investor's belief:**

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

- ▶ **AI's belief** depends on **memory**.

AI EFFECT

Sam Altman on GPT-6: 'People want memory'

PUBLISHED TUE, AUG 19 2025-8:30 AM EDT | UPDATED TUE, AUG 19 2025-8:51 AM EDT

**MacKenzie Sigalos**
@KENZIESIGALOSSHARE [f](#) [X](#) [in](#) [✉](#)

KEY POINTS

- Sam Altman said he wants future ChatGPT versions to let users define tone and personality.
- Altman called enhanced memory his favorite feature this year.
- "People want product features that require us to be able to understand them," Altman said last week.



OpenAI CEO Sam Altman, pictured, speaks with SoftBank Group CEO Masayoshi Son at an event in Tokyo on Feb. 3, 2025.

AI Memory κ

- ▶ Cost: Transformer with quadratic difficulty scaling
- ▶ Development: built-in session memory, context window.

AI Memory κ

- ▶ Cost: Transformer with quadratic difficulty scaling
- ▶ Development: built-in session memory, context window.
- ▶ LLM absorbs each signal d_{s_t} w.p. $\kappa \in [0, 1]$.

AI Memory κ

- ▶ Cost: Transformer with quadratic difficulty scaling
- ▶ Development: built-in session memory, context window.
- ▶ LLM absorbs each signal ds_t w.p. $\kappa \in [0, 1]$.
- ▶ Beliefs in $z_t \equiv \ln \frac{p_t}{1-p_t}$ and $\hat{z}_t \equiv \ln \frac{\hat{p}_t}{1-\hat{p}_t}$:

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

AI Memory κ

- ▶ Cost: Transformer with quadratic difficulty scaling
- ▶ Development: built-in session memory, context window.
- ▶ LLM absorbs each signal ds_t w.p. $\kappa \in [0, 1]$.
- ▶ Beliefs in $z_t \equiv \ln \frac{p_t}{1-p_t}$ and $\hat{z}_t \equiv \ln \frac{\hat{p}_t}{1-\hat{p}_t}$:

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

- ▶ **No memory** ($\kappa = 0$)
 - ▶ $\hat{p}_t = \hat{p}_0$.
- ▶ **Full memory** ($\kappa = 1$) + aligned prior ($\hat{p}_0 = p_0$):
 - ▶ $\hat{p}_t = p_t$ (AI tracks investor perfectly)

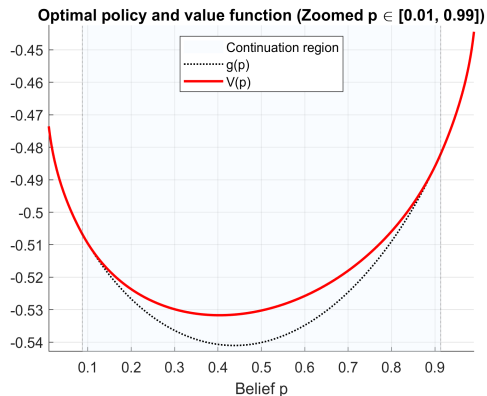
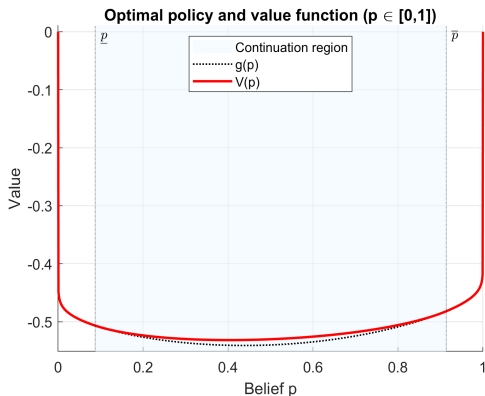
Value Function and Optimal Policy

- ▶ **Threshold policy:** continue iff $p \in (\underline{p}, \bar{p})$.
- ▶ **Continuation region HJB:**

$$\underbrace{0}_{\text{req. return}} = \underbrace{-c}_{\text{flow cost}} + \underbrace{\lambda(g(p) - V(p))}_{\text{Poisson shock}} + \underbrace{\frac{p^2(1-p)^2}{2\sigma^2} V''(p)}_{\text{learning}}$$

Eqm and Implications

Equilibrium: $\kappa > 0$



- $V(p)$ solved in closed form.
- $\kappa = 0$; $\kappa = 1 \& p_0 = \hat{p}_0$: closed form up to one unknown.

When does Memory Help?

1. Prompt quality substitutes for memory

- ▶ High prompt quality $\Rightarrow$ belief converges to certainty
- ▶ Even minimal AI memory achieves first best
- ▶ Good prompts reduce need for long context windows

When does Memory Help?

1. Prompt quality substitutes for memory

- ▶ High prompt quality $\Rightarrow$ belief converges to certainty
- ▶ Even minimal AI memory achieves first best
- ▶ Good prompts reduce need for long context windows

2. Memory benefits depends on prior alignment $(p_0, \hat{p}_0)$

- ▶ Aligned priors: more memory always helps

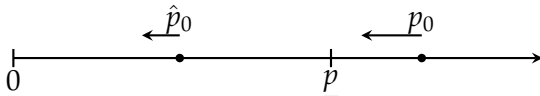
When does Memory Help?

1. Prompt quality substitutes for memory

- ▶ High prompt quality $\Rightarrow$ belief converges to certainty
- ▶ Even minimal AI memory achieves first best
- ▶ Good prompts reduce need for long context windows

2. Memory benefits depends on prior alignment $(p_0, \hat{p}_0)$

- ▶ Aligned priors: more memory always helps
- ▶ Misaligned priors: more memory can *hurt* bc overshooting



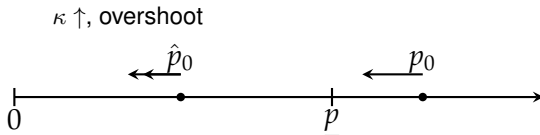
When does Memory Help?

1. Prompt quality substitutes for memory

- ▶ High prompt quality $\Rightarrow$ belief converges to certainty
- ▶ Even minimal AI memory achieves first best
- ▶ Good prompts reduce need for long context windows

2. Memory benefits depends on prior alignment $(p_0, \hat{p}_0)$

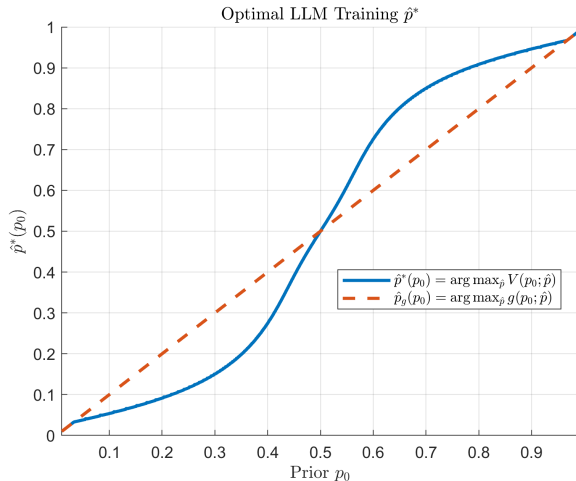
- ▶ Aligned priors: more memory always helps
- ▶ Misaligned priors: more memory can *hurt* bc overshooting



Optimal AI Training: “Opinionated” LLM

- ▶ If immediately stops, $\hat{p}_0^* = p_0$.
- ▶ Optimal AI prior $\hat{p}_0^*(\kappa = 0)$:

$$\hat{p}_0^* \begin{cases} > p_0 & \text{if } p_0 > 0.5 \\ < p_0 & \text{if } p_0 < 0.5 \\ = p_0 & \text{if } p_0 = 0.5 \end{cases}$$



Optimal AI Training: “Opinionated” LLM

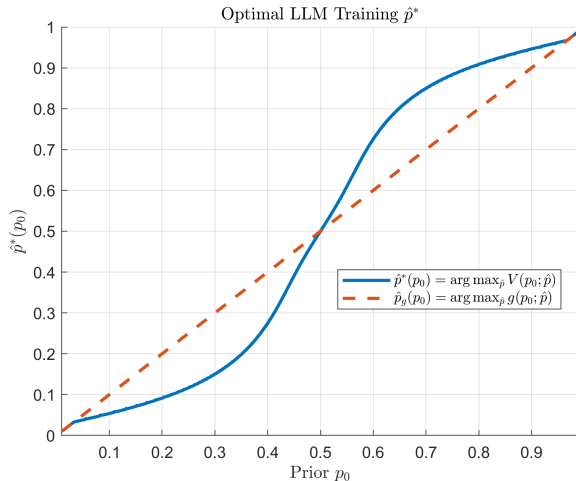
► If immediately stops, $\hat{p}_0^* = p_0$.

► Optimal AI prior $\hat{p}_0^*(\kappa = 0)$:

$$\hat{p}_0^* \begin{cases} > p_0 & \text{if } p_0 > 0.5 \\ < p_0 & \text{if } p_0 < 0.5 \\ = p_0 & \text{if } p_0 = 0.5 \end{cases}$$

► **Option value of learning**

► Informativeness: favors extreme $\hat{p}$



Empirical Simulations

Simulation Design

- ▶ **Data:** 500 investor profiles from SCF 2022
 - ▶ Ground truth: Vanguard 11-question questionnaire (preference ω)
 - ▶ Vanguard recommendation ($a^{FB} = \omega\theta_1 + (1 - \omega)\theta_0$)

Simulation Design

- ▶ **Data:** 500 investor profiles from SCF 2022
 - ▶ Ground truth: Vanguard 11-question questionnaire (preference ω)
 - ▶ Vanguard recommendation ($a^{FB} = \omega\theta_1 + (1 - \omega)\theta_0$)
- ▶ **LLM:** OpenAI GPT-5, two-sided simulation (advisor $\hat{p}$, investor p)
 - ▶ Specify prior (p_0) $\rightarrow$ Q&A loop (ds_t) $\rightarrow$ stop (Poisson or $\underline{p}, \bar{p}$) $\rightarrow$ rec (m^L) $\rightarrow$ action (a)

Simulation Design

- ▶ **Data:** 500 investor profiles from SCF 2022
 - ▶ Ground truth: Vanguard 11-question questionnaire (preference ω)
 - ▶ Vanguard recommendation ($a^{FB} = \omega\theta_1 + (1 - \omega)\theta_0$)
- ▶ **LLM:** OpenAI GPT-5, two-sided simulation (advisor $\hat{p}$, investor p)
 - ▶ Specify prior (p_0) $\rightarrow$ Q&A loop (ds_t) $\rightarrow$ stop (Poisson or $\underline{p}, \bar{p}$) $\rightarrow$ **rec** (m^L) $\rightarrow$ action (a)
- ▶ **What is fed to the advisor (memory κ):**
 - ▶ No memory: last Q&A pair ($\hat{p}_t = \mathbb{E}[\omega|ds_{t-}]$, $\kappa = 0$)
 - ▶ Full memory: entire chat history ($\kappa = 1$)
 - ▶ Full information: complete questionnaire ($\hat{p} = \omega$)

H1: Self-Learning Dominates

Stage	Accuracy
Prior (before any interaction)	71.3
Final Before Rec (after Q&A, before rec)	85.4
Final (after seeing recommendation)	87.0
Improvement from interaction alone	+14.1pp
Additional improvement from recommendation	+1.6pp

Finding: Primary value is investor **self-learning**. Each additional Q&A round adds ~ 1pp to accuracy.

H2: Impatience Hurts

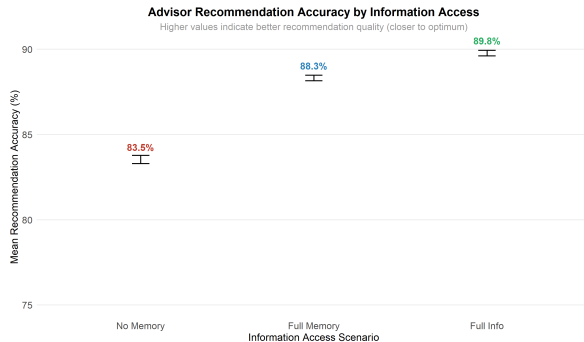
	Accuracy _i		
	(1)	(2)	(3)
Early Termination	−2.62*** (0.45)	−0.96** (0.42)	−1.12** (0.44)
Controls	NO	YES	YES

Finding: Exogenous termination reduces accuracy even controlling for rounds/words $\Rightarrow$ validates that **impatience hurts**.

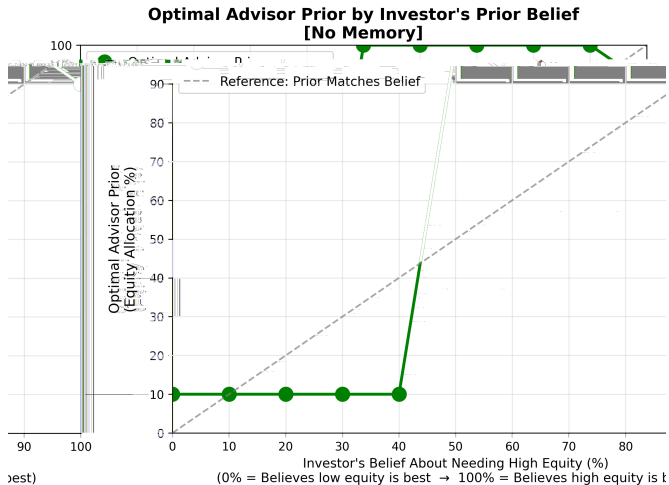
H3: Memory Helps

Condition	Accuracy
No memory	83.5
Full memory	88.3
Full information	89.8

- Memory: +4.8pp
- Full info gap: +1.5pp



H4: Opinionated AI Optimal



Conclusion

- ▶ **Debiasing friction**—AI doesn't know user's preference.
- ▶ **Key insights:**
 1. Prompt quality can substitute for memory
 2. Memory benefits depend on prior alignment
 3. **Opinionated AI can be optimal**
- ▶ **Empirical validation:**
 - ▶ Self-learning dominates (+14.1pp)
 - ▶ Memory helps (+4.8pp)
 - ▶ Opinionated AI pattern validated