

NBER WORKING PAPER SERIES

SEEING THE GOAL, MISSING THE TRUTH:
HUMAN ACCOUNTABILITY FOR AI BIAS

Sean S. Cao
Wei Jiang
Hui Xu

Working Paper 35142
<http://www.nber.org/papers/w35142>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
April 2026

We thank seminar participants at the University of Connecticut, the University of Maryland, Lancaster University, the University of California-Riverside, Hitotsubashi University, Waseda University, Singapore Management University, Nanyang Technological University, Tulane University, Cheung Kong Graduate School of Business (CKGSB), Peking University, the University of Central Florida, Baruch College, Stony Brook University, UBS, and the China International Conference in Finance (CICF) for helpful feedback. Sean Cao gratefully acknowledges support from the Smith AI Initiative for Capital Markets Research at the University of Maryland. The views expressed herein are those of the authors and do not necessarily reflect the views of the National Bureau of Economic Research.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2026 by Sean S. Cao, Wei Jiang, and Hui Xu. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

Seeing the Goal, Missing the Truth: Human Accountability for AI Bias
Sean S. Cao, Wei Jiang, and Hui Xu
NBER Working Paper No. 35142
April 2026
JEL No. G14, G17

ABSTRACT

This research explores how human-defined goals influence the behavior of Large Language Models (LLMs) through purpose-conditioned cognition. Using financial prediction tasks, we show that revealing the downstream use (e.g., predicting stock returns or earnings) of LLM outputs leads the LLM to generate biased sentiment and competition measures, even though these measures are intended to be downstream task-independent. Goal-aware prompting shifts these intermediate measures toward the disclosed downstream objective, producing in-sample overfitting. Specifically, purpose leakage improves performance on data prior to the LLM’s knowledge cutoff, but provides no advantage after the cutoff. This bias is strong enough that regularization of prompt instructions cannot fully address this form of overfitting. We further show that the bias can arise from users’ unintentional conversational context that hints at the purpose. Overall, we document that AI bias due to “seeing the goal” is not an algorithmic flaw, but stems from human accountability in research design.

Sean S. Cao
University of Maryland, College Park
scao824@umd.edu

Hui Xu
Lancaster University
h.xu10@lancaster.ac.uk

Wei Jiang
Emory University
Goizueta Business School
and NBER
wei.jiang@emory.edu

Seeing the Goal, Missing the Truth: Human Accountability for AI Bias

ABSTRACT

This research explores how human-defined goals influence the behavior of Large Language Models (LLMs) through purpose-conditioned cognition. Using financial prediction tasks, we show that revealing the downstream use (e.g., predicting stock returns or earnings) of LLM outputs leads the LLM to generate biased sentiment and competition measures, even though these measures are intended to be downstream task-independent. Goal-aware prompting shifts these intermediate measures toward the disclosed downstream objective, producing in-sample overfitting. Specifically, purpose leakage improves performance on data prior to the LLM’s knowledge cutoff, but provides no advantage after the cutoff. This bias is strong enough that regularization of prompt instructions cannot fully address this form of overfitting. We further show that the bias can arise from users’ unintentional conversational context that hints at the purpose. Overall, we document that AI bias due to “seeing the goal” is not an algorithmic flaw, but stems from human accountability in research design.

1 Introduction

In organizational settings, the disclosure of downstream use can alter the nature of an intermediate task. Consider a human assistant asked to summarize interview transcripts post hoc. If told that the summaries will be used to evaluate recruiting effectiveness, the assistant may emphasize the strengths and “bright spots” of candidates who ultimately receive offers, while downplaying uncertainty, mixed signals, or unfavorable assessments. The resulting summaries support a narrative of successful recruitment, yet are less faithful to the underlying interviews. This behavioral shift reflects rational adaptation to perceived downstream purpose rather than error or lack of effort. More broadly, when intermediate tasks are conditioned on their eventual application, neutrality can be compromised even in the absence of explicit incentives.

Purpose-conditioned cognition is a well-documented human trait.¹ Individuals often adjust interpretation and judgment in response to perceived goals or downstream use, even without explicit incentives; and such behavior is difficult to eliminate. As a result, human-generated intermediate outputs are understood to embed goal-related distortions. On the other hand, one

¹Related ideas appear in the literature on motivated reasoning and belief distortion (Bénabou and Tirole, 2016), framing effects in decision-making (Tversky and Kahneman, 1981) and the limits of attention and interpretation in complex environments (Mullainathan et al., 2008). While the terminology differs across fields, these studies share the insight that cognition and interpretation are systematically shaped by perceived objectives and contextual framing.

might expect algorithmic systems, including large language models (LLMs), to be immune from human biases. This presumption motivates a natural question: does purpose-conditioned behavior also arise in LLMs? In other words, if we restrict AI to “assistant” roles, does this separation from the final decision task ensure neutral and unbiased intermediate outputs?

Our research design addresses this question by holding the input text, model, and scoring task fixed, while varying only whether the downstream use of the output is disclosed. Specifically, we prompt two large language models (ChatGPT and Gemini) to generate numerical measures of sentiment and competition intensity from the same source information (earnings call transcripts) but under two prompt conditions. In the goal-blind condition, the model is asked to produce a score without any reference to its eventual use. In the goal-aware condition, the model is informed that its output will be used to predict future stock returns or earnings. This design therefore isolates the effect of disclosing downstream use on the model’s intermediate outputs.

Disclosing downstream use materially alters both the statistical properties and the economic content of LLM-generated sentiment scores. In standard portfolio-sorting tests, goal-aware sentiment produces substantially larger return spreads than goal-blind sentiment prior to the model’s knowledge cutoff. This pattern suggests that the model conditions its intermediate outputs on the disclosed objective and implicitly re-weights return-relevant patterns learned during training. The relative advantage of the goal-aware regime disappears after the knowledge cutoff, indicating that the earlier performance gap is tied to information availability rather than a stable improvement in signal quality.

This comparison of the two prompt regimes is further confirmed in [Fama and MacBeth \(1973\)](#) return-prediction regressions. Intermediate variables generated under goal-aware prompts exhibit stronger incremental predictive power than those from the goal-blind regime in the pre-cutoff period, both in coefficient significance and in out-of-sample R^2 . Once the model no longer has access to future information beyond the cutoff, this advantage dissipates. Together, these results are consistent with objective-conditioned output adjustment rather than a structurally more informative sentiment measure. In fact, the enhanced goodness-of-fit prior to knowledge cutoff even slightly degrades out-of-sample model generalization. On the surface, the pre-cutoff inflation and post-cutoff attenuation of goal-aware advantage look like the p-hacking and data-mining problems

documented in the anomalies literature.² But the mechanism is different. P-hacking distorts because researchers cycle through specifications and report the ones that survive conventional thresholds. In our setting, the inflation occurs within a single forward pass. The model shifts its outputs toward return-predictive patterns as soon as it is told what those outputs are for. Pre-registration and multiple-testing corrections cannot address this, because there is only one test.

One might expect that a carefully worded instruction can undo what disclosure triggers, but we further show that explicit disclosure is not required. “Subtle” clues (e.g., asking the LLM whether sentiment scores predict stock returns, with no instruction linking that context to the scoring task) generate largely the same outcome. Even an explicit directive asking the model to minimize any reliance on outcomes beyond the measurement window, effectively auditing its own reasoning for hindsight, merely attenuates the bias instead of eliminating it. Once information is encoded in model weights, prompt-level instructions cannot selectively suppress it at inference time.

The seeming goal-conditioning behavior of a cold-blooded algorithm has its roots in how AI systems are trained and deployed. Large language models are optimized to generate outputs that satisfy the objective implied by the prompt, conditional on their learned representations. Disclosing downstream use alters this implicit objective, encouraging responses that align with anticipated evaluation criteria rather than with neutral processing of the input. Put differently, while a goal-blind prompt positions the LLM as a neutral sensor, a goal-aware prompt transforms it into a strategic tuner. This behavior does not require intent or explicit embedding of forward-looking information. Instead, it reflects the model’s reliance on correlations learned during training and its sensitivity to contextual cues about what constitutes a desirable output. As a result, goal awareness can improve in-sample alignment while reducing robustness when the information environment changes.

Our study contributes to the AI and LLM methodology literature (reviewed in Appendix A) in two distinct ways. First, existing research largely attributes AI bias to algorithmic limitations. Look-ahead bias and memorization, for example, are commonly traced to models’ access to unintended training data (Glasserman and Lin, 2023; Sarkar and Vafa, 2024; Lopez-Lira, Tang and Zhu, 2025; He, Lv, Manela and Wu, 2025; Cao, Wang and Xiang, 2025). Moving beyond these training data-centric issues, AI has also been criticized for potential biases in exploration strategies and model

²See, e.g., Harvey et al. (2016); Hou et al. (2020); McLean and Pontiff (2016)

weights, which may be related to ethical concerns (Fedyk et al., 2024) and preference alignment (Ouyang et al., 2024; Hirshleifer et al., 2025; Bini et al., 2026). Our analysis shifts the focus from these machine-level limitations to human use of machine. We show that human disclosure of downstream task (e.g., return and earnings predictions) reshapes intermediate outputs (e.g., sentiment and competition measures) in ways that inflate the downstream task performance. Ultimately, this leads to impressive in-sample results that fail to generalize out-of-sample.

The second contribution of the paper is to isolate goal awareness from other AI biases, which naturally emerge when LLMs are deployed directly on end tasks such as return forecasting. A commonly proposed safeguard is to limit LLMs to intermediate functions, analogous to a research assistant, while reserving final predictions and decisions for human judgment. Our results show that even at the intermediate-measurement stage, disclosure of the downstream objective can introduce additional, goal-conditioned distortion in the generated signals. The conversation history shows that a model with encountered relevant context infers the downstream objective and adjusts without any direct instruction. Prompt-level regularization can reduce but not fully remove this bias, since the underlying knowledge is distributed across model parameters rather than stored in a form that instructions can selectively override.

More broadly, the analysis clarifies the distinction between using AI as a task agent that optimizes directly toward an announced objective and using it as a measurement tool that produces inputs for subsequent human evaluation. Prompt and workflow designs that deliberately reduce objective conditioning can therefore mitigate machine-driven bias. The mechanism we uncover thus bears an analogy to the phenomenon of “AI sycophancy” discussed in recent work, e.g., Sharma et al. (2023). In both cases, large language models adjust their outputs in response to contextual cues about what constitutes a desirable response, rather than adhering to a task-invariant notion of correctness or neutrality.³ The underlying commonality is not intent, but sensitivity to signals about reward or approval embedded in the prompt. As a result, both phenomena illustrate how models trained to be helpful and aligned can deviate from neutral information processing when the prompt conveys implicit incentives, even in the absence of explicit instructions or forward-looking information.⁴ A design choice that weakens alignment with a stated downstream goal can improve

³In sycophancy, this adjustment manifests as alignment with a user’s stated beliefs or preferences; in our setting, it appears as alignment with an inferred evaluation objective.

⁴In computer science, related concerns are discussed under reward hacking, specification gaming, and objective

statistical validity and out-of-sample reliability.

In a nutshell, intermediate variables constructed with AI assistance should, whenever feasible, be generated under prompts agnostic to downstream use and evaluated with strict out-of-sample tests. Treating LLMs as neutral measurement devices requires not only holding inputs and models fixed, but also constraining the informational environment to limit objective inference by the model beyond the intrinsic requirements of the task. In this sense, the challenge is not primarily one of algorithmic bias, but of human accountability in how objectives, context, and evaluation criteria are built in the system.

2 Experimental Design and Data

2.1 LLM Prompt and Scoring

Our objective is to examine whether large language models (LLMs) systematically adjust their outputs when they are informed about the downstream task for which those outputs will ultimately be used. We focus on two economically important forecasting applications: monthly stock returns and quarterly earnings per share (EPS). In both settings, the firm’s most recent earnings call transcript serves as the sole input to the LLM. Rather than asking the model to directly predict economic outcomes, we instruct it to generate an intermediate score that is subsequently used in a simple predictive regression.

- Return prediction task. For each calendar month t , the model is prompted with the most recent earnings call transcript available at $t - 1$ to generate a continuous sentiment score, ranging from -1 to 1, summarizing the firm’s business sentiment for month t . We then use this month- t sentiment score to predict the firm’s stock returns in month t .
- Earnings prediction task. For each fiscal quarter t , the model is prompted with the most recent earnings call transcript available at $t - 1$ to generate a competition score, ranging from -1 to 1, which captures the intensity of competitive pressure the firm faces in quarter t . We then use this quarter- t competition score to predict the firm’s earnings realized in quarter t .

misgeneralization, where models optimize inferred proxy objectives rather than the intended task. Our setting requires no change in model parameters or rewards; the distortion arises solely from human framing at inference time.

A central feature of the experimental design is that the LLM does not directly forecast stock returns or earnings. Instead, it produces an *intermediate* construct (i.e., sentiment or competitive intensity), which we subsequently map to economic outcomes using standard econometric models. This design choice allows us to test whether the LLM adjusts these intermediate scores as if it were strategically responding to the knowledge of their downstream use. For each task, we construct two prompts that are otherwise identical in wording and structure, differing only in whether the prompt explicitly discloses the ultimate use of the generated score.

1. Goal-Blind Prompt (control): The LLM is asked to generate the score without being told that it will be used in a forecasting regression.
2. Goal-Aware Prompt (treatment): The LLM is informed that the score it produces will later be used as an explanatory variable in a regression to predict either stock returns or earnings.

This minimalist treatment variation allows us to cleanly isolate whether, and through what channels, an LLM alters its outputs once it becomes aware of the user’s ultimate objective. Conceptual concerns arise from anticipatory behavior, in which the model conditions its responses on inferred downstream use rather than on the stated prediction task alone; reward hacking, whereby the model optimizes for the perceived evaluation criterion instead of the intended informational objective; and goal misalignment, in which the model’s internal optimization departs from the user’s declared constraints. Importantly, these issues are examined within an economically meaningful forecasting environment, allowing deviations in model behavior to be traced to distortions in information processing rather than an abstract alignment failure.

To illustrate, the following text shows the goal-blind prompt used to generate sentiment scores:

```
"For the following tasks, all dates are expressed in the format
MM/DD/YYYY (month/day/year) .

Below is the earnings call transcript of {ticker}. Please provide a
continuous sentiment score in [-1, 1] about the firm's business
sentiment for the month ending on {date}.

Provide a precise numerical answer. Format as a JSON object with the
following fields:

- answer: The precise numerical answer to the question. No strings.
{the firm's earnings call transcript} ."
```

The corresponding goal-aware version differs only by the addition of a single sentence that discloses the downstream use of the output, namely the predictive task.

```
"For the following tasks, all dates are expressed in the format
MM/DD/YYYY (month/day/year).

Below is the earnings call transcript of {ticker}. Please provide a
continuous sentiment score in [-1, 1] about the firm's business
sentiment for the month ending on {date}. The sentiment score later will be
used as an explanatory variable in a regression to predict the monthly stock returns ending
on {date}.

Provide a precise numerical answer. Format as a JSON object with the
following fields:

- answer: The precise numerical answer to the question. No strings.
{the firm's earnings call transcript}."
```

Aside from this one sentence, all instructions, including input data, output format, and numerical constraints, remain identical across the two prompts and processes.⁵

2.2 Data and LLM models

Our sample consists of S&P 500 firms over the period from January 2022 to December 2024. Earnings call transcripts are obtained from Capital IQ, which provides standardized, time-stamped transcripts for publicly listed firms. Monthly stock return data are drawn from CRSP, and accounting information, including earnings per share (EPS), is sourced from Compustat. We further use CRSP and Compustat data to construct standard firm-level control variables employed in our predictive regressions, such as the book-to-market ratio and firm size.

To generate sentiment and competition scores, our baseline experiment uses the GPT-4o-mini model. GPT-4o-mini ("o" for "omni") is a lightweight, cost-efficient language model designed for focused inference tasks with low latency. According to OpenAI, GPT-4o-mini is produced via distillation from a larger frontier model (GPT-4o), allowing it to replicate much of the larger model's behavior at substantially lower computational cost. Critically for our experimental design, GPT-4o-mini has a fixed knowledge cutoff of October 1, 2023. As a result, the model does not

⁵The full prompts are also detailed in Appendix B.

have access to information that occurs after this date. This feature is central to our analysis, as it allows us to cleanly separate changes in predictive performance driven by prompt design and goal awareness from those arising from direct access to future information.

As a robustness test, we repeat the same analyses using Gemini 2.5, a large language model developed by Google with a well-documented knowledge cutoff date of January 2025.⁶

2.3 Evaluation Metrics

We first evaluate the economic relevance of GPT-generated sentiment scores using portfolio-sorting tests that are standard in the asset-pricing literature. In each period, firms are sorted into quintiles based on their GPT-produced sentiment scores, and we form an equally weighted zero-investment long–short portfolio that buys firms in the highest quintile and sells firms in the lowest quintile. Portfolio performance is measured using average excess returns. The analysis is conducted separately for goal-aware and goal-blind scores. Comparing portfolio performance across the two prompt designs allows us to assess whether revealing the ultimate prediction task leads to economically stronger trading signals, and whether the relative performance of goal-aware versus goal-blind scores changes after the cutoff date.

We evaluate predictive performance using two complementary approaches: [Fama and MacBeth \(1973\)](#) predictive regressions and genuine out-of-sample forecasting. The first approach tests whether LLM-generated scores significantly predict returns or earnings in the full sample and how this relationship changes around the model’s knowledge cutoff. The second assesses practical forecasting accuracy by simulating real-time predictions on unseen data. Together, these methods provide a robust validation: the regressions establish statistical significance of the predictive variable, while the out-of-sample exercise determines whether this significance translates into reliable practical performance.

We use standard cross-sectional predictive regressions to examine statistical predictability. For stock returns, we implement the [Fama and MacBeth \(1973\)](#) methodology and estimate monthly

⁶Google does not disclose an exact knowledge cutoff date for Gemini 2.5, indicating only January 2025. We therefore set the cutoff to January 1, 2025. Given that our downstream prediction task is based on monthly returns measured at the last trading day of each month, this convention is unlikely to materially affect our results.

cross-sectional regressions of the form:

$$R_{i,t} = \beta_{1,t} \text{Score}_{i,t-1} \times \text{Pre-Cutoff}_t + \beta_{2,t} \text{Score}_{i,t-1} \times \text{Post-Cutoff}_t + \beta_{3,t} \text{Diff}_{i,t-1} \times \text{Pre-Cutoff}_t + \beta_{4,t} \text{Diff}_{i,t-1} \times \text{Post-Cutoff}_t + \alpha_t + \epsilon_{i,t}, \quad (1)$$

where $R_{i,t}$ denotes the excess return of firm i in month t . $\text{Score}_{i,t-1}$ is the sentiment score generated by the goal-blind GPT prompt using the most recent earnings call transcript available at $t - 1$. $\text{Diff}_{i,t-1}$ is the difference between the sentiment scores produced by the goal-aware and goal-blind prompts. Because the sentiment score does not have a natural scale, raw values may not be directly comparable across firms, industries, or time periods. For this reason we transform both sentiment scores into percentiles across all firms within each year-month so that $\text{Diff}_{i,t-1}$ variable captures the gap in the percentile values of scores generated under goal-aware and goal-blind scores. The time-period fixed effect α_t is included to absorb aggregate time series shifts in returns.

We interact both $\text{Score}_{i,t-1}$ and $\text{Diff}_{i,t-1}$ with indicator variables Pre-Cutoff_t and Post-Cutoff_t , which equal one if month t occurs before or after the LLM’s knowledge cutoff date, respectively. This specification allows predictive coefficients to differ across the two periods and is designed to isolate the role of goal awareness. When the LLM is goal aware, it may leverage patterns, associations, and correlations embedded in the full training sample available that are correlated with future outcomes. Because such information is internalized during training and cannot be selectively “unlearned” at inference time, goal awareness can thus induce outputs that are implicitly forward-looking relative to the evaluation period, leading to stronger predictive performance prior to the cutoff. Such behavior does not require the LLM to explicitly access forward-looking information in making predictions. Once the evaluation period extends beyond the model’s knowledge cutoff, however, these internalized correlations become equally stale, and the performance advantage of goal-aware scores over goal-blind scores should dissipate or even reverse.

Accordingly, any differential predictive power attributable to goal awareness should be concentrated in the pre-cutoff period and attenuated in the post-cutoff period. For this reason, we expect $\beta_{1,t}$ and $\beta_{2,t}$ in Equation (1) to be positive if the intermediate variable is informative. Moreover, the two coefficients should be of similar magnitude, as knowledge of future returns prior to the cutoff date should not matter under goal blindness. By the same reasoning, $\beta_{4,t} = 0$ since the lack

of future-relevant information renders any advantage of goal awareness nil after the cutoff. Finally, $\beta_{3,t}$ is the key coefficient of interest. Under the null, $\beta_{3,t} = 0$ if goal awareness does not change the way the LLM generates outputs; under the alternative, $\beta_{3,t}$ may be positive if disclosure of the downstream use of the output motivates the LLM to produce intermediate measures that better align with the ultimate task.

For earnings outcomes, we adopt an analogous specification, replacing monthly excess returns with quarterly earnings per share (EPS). We use diluted EPS excluding extraordinary items from Compustat, following common practice in the literature. In this setting, $Score_{i,t-1}$ corresponds to the competition score generated by the goal-blind prompt, transformed into percentiles within the same 4-digit SIC industry $\times$ year cohort. The difference measure captures the gap between the scores generated under the goal-blind and goal-aware prompts. This structure, which is parallel to our earlier exercise on stock returns, facilitates multiple validation by allowing us to test how goal awareness affects predictability across both financial and accounting settings.

A complementary approach to the cross-sectional predictive regression is assessing genuine, out-of-sample forecasting performance. At the beginning of each forecast period, T , we estimate the regression model on all data available through $T - 1$ and generate predictions for outcomes at T ; this estimation window is then expanded recursively to include additional observations before each new forecast, ensuring that every prediction is made using the maximum available history.⁷ The unit of observation is monthly for stock returns and quarterly for earnings. For stock returns prediction, this amounts to the following predictive regression:

$$R_{i,t} = \alpha + \gamma \text{Score}_{i,t-1} + \epsilon_{i,t}, \forall t \leq T - 1. \quad (2)$$

We then use the estimated coefficients $\{\hat{\alpha}, \hat{\gamma}\}$ to predict stock returns $R_{i,T}$ based on sentiment scores observed at $T - 1$. We estimate this regression separately using scores generated from the goal-aware and goal-blind prompts. Earnings predictions follow an analogous specification, with one modification: rather than using the level of EPS directly, we replace it with the year-over-year growth rate of same-quarter EPS, defined as $\frac{EPS_{i,T}}{EPS_{i,T-4}}$. This adjustment accounts for the

⁷This procedure is equivalent to supervised learning with an expanding training window: at each forecast period T , the model is trained on all observations $\{1, \dots, T - 1\}$ and evaluated on the held-out observation at T . The training set grows monotonically with each step, and no future data ever enters the training set prior to the corresponding forecast, a design that strictly preserves temporal holdout and precludes look-ahead bias.

well-documented seasonality and serial correlation in earnings.

Once we obtain model-based forecasts, we construct a forecast accuracy measure analogous to the out-of-sample (OOS) R^2 (Campbell and Thompson, 2008):

$$R_{OOS,i,T}^2 = 1 - \frac{(y_{i,T} - \hat{y}_{i,T})^2}{(\bar{y}_{T-1} - y_{i,T})^2}, \quad (3)$$

where $y_{i,T}$ is the realized outcome, $\hat{y}_{i,T}$ is the model-based forecast, and $\bar{y}_{T-1}$ is the benchmark forecast—specifically, the simple historical cross-sectional mean of the outcome variable across all firms, calculated using data available through period $T - 1$.

With the forecast accuracy measure at the firm-period level, we are able to assess their performance in relation to goal-awareness and interacted with the LLM knowledge cutoff date. More specifically, we estimate the following panel regression:

$$R_{OOS,i,T}^2 = \theta_1 \text{Goal-Aware}_{i,T} \times \text{Post-Cutoff}_T + \theta_2 \text{Goal-Aware}_{i,T} + \theta_3 \text{Post-Cutoff}_T + \mu_i + \nu_T + \epsilon_{i,T} \quad (4)$$

where $\text{Goal-Aware}_{i,T}$ is an indicator equal to one (zero) if $R_{OOS,i,T}^2$ results from the goal-aware (goal-blind) regime. μ_i and ν_T denote firm and time fixed effects, respectively. The coefficient θ_1 is of central interest, as it captures the change in incremental forecasting performance of goal-aware relative to goal-blind scores before and after the model’s knowledge cutoff date. A negative θ_1 indicates that the incremental forecasting performance of goal-aware prompt weakens after the knowledge cutoff.

3 Findings

3.1 Sentiment Scores and Stock Returns Prediction

3.1.1 Portfolio Sorts and Return Performance

We use GPT as the primary LLM model for illustration, followed by replicated findings on Gemini in Section 5.3. We begin by examining the economic implications of GPT-generated sentiment scores using portfolio-sorting tests. The purpose of this exercise is to verify that the sentiment scores have

predictive content for stock returns. Although return forecasting is not a central objective of this study, the presence of a predictive variable is required to evaluate how the LLM’s output changes when the model is made aware of the downstream predictive task. Figure 1 plots the cumulative returns to long–short portfolios formed by buying firms in the highest sentiment quintile and selling firms in the lowest sentiment quintile, separately for scores generated under goal-aware and goal-blind prompts. Table 1 reports the corresponding average monthly return spreads for the pre- and post-knowledge cutoff periods.

Prior to the knowledge cutoff, sentiment scores generated under both prompt designs exhibit economically meaningful return predictability. As reported in Table 1, the long–short portfolio based on goal-aware sentiment earns an average monthly return spread of 1.552%, while the corresponding spread based on goal-blind sentiment is 1.069%. Both spreads are statistically significant, indicating that GPT-generated sentiment contains predictive information about future stock returns even when the ultimate objective is not fully communicated to the model. Notably, the goal-aware strategy delivers significantly stronger performance: the difference in High–Low spreads between the two regimes, 0.483 percentage point per month, is statistically significant at the 5% level. This relative outperformance is evident graphically in Figure 1, where cumulative returns of the goal-aware portfolio diverge upward from those of the goal-blind portfolio in the period leading up to the cutoff.

After the knowledge cutoff, both strategies continue to generate statistically significant return spreads. The monthly High–Low spread equals 2.269% for the goal-aware portfolio and 2.239% for the goal-blind portfolio, with no economically or statistically meaningful difference between the two. Confirming this result, the cumulative return paths in the right panel of Figure 1, normalizing both portfolios to start from zero at the cutoff date, track each other closely in the post-cutoff period, showing no persistent advantage for the goal-aware strategy. If anything, the goal-blind long–short portfolio performs slightly better. One possible explanation is mild overfitting: making the LLM aware of the downstream objective may induce greater optimization to in-sample patterns, whereas the goal-blind specification generates more stable signals that generalize slightly better in the true out-of-sample (post-cutoff) period.

Taken together, the portfolio evidence shows a contrast around the knowledge cutoff. Before the cutoff, disclosing the downstream prediction task increases the measured economic content of

GPT-generated sentiment scores. After the cutoff, this relative difference disappears, even though both prompt designs continue to exhibit statistically significant return predictability in absolute terms. This pattern indicates that the higher pre-cutoff performance of goal-aware sentiment does not arise from superior information extraction. Instead, it is consistent with goal awareness reshaping the distribution of model outputs by incorporating back-tested performance when data from the evaluation period are available during training. The absence of a post-cutoff difference helps identify the source of the pre-cutoff effect. Overall, the results highlight the importance of accounting for goal awareness when interpreting LLM-based measures in empirical studies.

3.1.2 Predictive Regressions and Out-of-Sample Performance

The comparative predictive performance of the goal-aware and goal-blind regimes can also be examined within the regression framework described in Section 2.3. Table 2 reports Fama–MacBeth estimates from monthly cross-sectional regressions of excess stock returns on GPT-generated sentiment scores, following Equation (1). The specification allows the slope coefficients to differ between the goal-aware and goal-blind regimes and across the pre- and post-knowledge cutoff periods.

Across specifications, sentiment scores generated under the goal-blind prompt display economically meaningful and statistically significant predictive power both before and after the knowledge cutoff. The estimated coefficients on the goal-blind score interacted with the pre-cutoff indicator are positive and statistically significant, and their magnitudes remain similar in the post-cutoff period. Formal tests reported in the bottom panel do not reject equality between the pre- and post-cutoff coefficients for the goal-blind score.

By contrast, the incremental predictive content of goal-aware sentiment—captured by the *Diff* measure, exhibits a discontinuity at the knowledge cutoff. In the pre-cutoff period, the coefficient on *Diff* is positive and statistically significant across both specifications, indicating that goal-aware sentiment contains incremental predictive power. The magnitude holds across firm characteristics such as size and book-to-market. After the knowledge cutoff, the *Diff* coefficient collapses to near-zero. Consistent with this pattern, the bottom panel reports that the equality of the pre- and post-cutoff *Diff* coefficients is rejected at the 5% significance level. Altogether, the regression

evidence closely mirrors the portfolio results.

Next we evaluate whether these regression-based differences carry over to out-of-sample forecasting performance as modeled in Equation (4). Figure 2 plots the monthly out-of-sample R^2_{OOS} obtained from forecasting next-month stock returns using sentiment scores generated under goal-aware and goal-blind prompts. In each month T , we estimate predictive regressions using an expanding window (where the estimation sample grows to include all available data up to the prior month $T - 1$), compute out-of-sample accuracy relative to a simple historical cross-sectional mean benchmark calculated using data available through the prior month $T - 1$, and average the resulting out-of-sample goodness-of-fit, $R^2_{OOS,i,T}$, across firms i in each time period T .

Figure 2 reveals a distinct shift in relative predictive accuracy across the knowledge cutoff. Prior to the cutoff, the goal-aware prompt generally outperforms the goal-blind prompt. Following the cutoff, however, this relationship reverses: the performance of the goal-aware prompt deteriorates significantly, falling below that of the goal-blind benchmark. Table 3 formalizes this comparison by regressing monthly R^2_{OOS} values on the interaction between the *Goal-aware* prompt and the *Post-Cutoff* indicators, following the specification in Equation (4). Consistent with the visual evidence, the coefficient on the interaction of the indicators is negative and statistically significant across specifications. The magnitude of the estimate implies a meaningful reduction in out-of-sample predictive accuracy for goal-aware sentiment once the model’s knowledge becomes stale.

The comparative results of out-of-sample predictive performance complement those from portfolio sorting and cross-sectional predictive regressions. While goal-aware sentiment appears to deliver stronger in-sample and pre-cutoff predictive signals, this advantage does not persist when evaluated under a strict out-of-sample framework after the knowledge cutoff. The pattern suggests that the superior pre-cutoff performance of goal-aware sentiment reflects, at least in part, the model’s tendency to condition its outputs on the evaluation objective in ways that do not generalize once the underlying knowledge environment changes.

Taken together, results from three complementary research designs indicate that making the downstream objective explicit can improve apparent predictive performance in-sample and prior to the knowledge cutoff, but does not improve, or may even degrade out-of-sample generalization. This pattern is consistent with goal awareness inducing objective-conditioned optimization by a LLM that exploits correlations present in the training or evaluation sample at the expense of

extracting stable predictive structure. The findings highlight the distinction between economically meaningful signals and prompt-induced optimization artifacts when employing LLM-generated measures in empirical research.

3.2 Earnings Prediction with Competition Scores

Similar to our discussion of GPT-generated sentiment scores and their relationship with future stock returns, we first study the ability of competition-threat scores to predict future earnings. Table 4 reports Fama–MacBeth estimates from firm-quarter panel regressions of next-quarter earnings per share on GPT-generated competition scores, allowing the predictive slopes to differ between goal-aware and goal-blind regimes and across the pre- and post-knowledge cutoff periods.

As in Table 2, we focus on *Diff*, the difference between goal-aware and goal-blind competition threat scores. Its interactions with *Pre-cutoff* and *Post-cutoff* capture how the goal-aware score’s incremental predictive power changes around the knowledge cutoff. Control variables follow So (2013): Column (1) includes the previous quarter’s truncated EPS (set to zero when negative) and a loss indicator for negative EPS; column (2) includes the full set of controls. The story that emerges is much the same as before: in the pre-cutoff period, the coefficient on *Diff* is negative and statistically significant at the 5% level across both specifications, consistent with the intuition that fiercer competition depresses earnings. This suggests that the goal-aware LLM tends to assign elevated competition threat scores to firms that subsequently report weaker earnings. Once we move past the knowledge cutoff, this effect vanishes and the *Diff* coefficients drop to near zero. Formal tests reported in the bottom panel reject the equality of the pre- and post-cutoff *Diff* coefficients at conventional significance levels.

We conclude our analysis by evaluating the out-of-sample forecasting power of GPT-generated competition scores for future earnings growth. Figure 3 plots quarterly average out-of-sample R^2_{OOS} values from expanding-window forecasts of earnings growth (where each forecast uses all available historical data up to the prior quarter) using both goal-aware and goal-blind competition scores. Table 5 complements this visual evidence with formal regression tests, which confirm that the incremental predictive content of goal-aware scores collapses after the knowledge cutoff.

In sum, the earnings predictions reinforce the central message: While GPT-generated competi-

tion scores exhibit economically intuitive relations with future earnings, the incremental predictive content attributable to goal awareness is confined to the pre-cutoff period. This confirms that communicating downstream purpose of the intermediate output may amplify predictive relations in-sample with no out-of-sample generalization.

4 Mechanism and Subtlety of Goal Awareness

4.1 A theoretical framework

The preceding evidence raises a natural question: why does informing the model of a downstream objective change its output at all? A simple framework provides an answer. Let the LLM’s output be an intermediate signal s intended to recover a latent sentiment θ , which predicts returns y . Under goal blindness, the model is a *measuring sensor*, minimizes a representation loss and outputs $s = \mathbb{E}[\theta \mid x]$, where x is input text. When told that s will be used to predict returns, the model is an *optimizing solver* and shifts toward a joint objective that trades off fidelity to sentiment against covariance with y .

This shift traces to how large language models are trained. LLMs are first pretrained to predict the next token given the full prompt and context. They are then fine-tuned via reinforcement learning from human feedback (RLHF), where annotators reward responses that are helpful and aligned with the user’s stated purpose. Because the reward model evaluates responses against the full prompt, contextual statements about intended use enter what the model treats as a “good” output. The model does not distinguish between an explicit optimization target and background context; both enter the effective objective through the reward function. Disclosure that the score will be used to predict returns therefore functions as part of the task definition, tilting s toward return-predictive patterns internalized during pretraining and reinforced through RLHF alignment. The resulting measure shows stronger in-sample predictability but fails to generalize once the knowledge environment changes. Goal awareness, in short, converts a neutral measurement tool into a partially optimized predictive proxy.

4.2 The subtlety of goal awareness

The obvious remedy seems simple enough: withhold any reference to downstream use and instruct the model to perform a narrowly defined intermediate task. In practice, this is harder than it sounds. Goal awareness arises not only from explicit instruction but from context, prior interactions, and learned associations, often in ways that are difficult to anticipate or control.

A parallel from human behavior is worth noting. Even when tasks are framed as neutral, individuals infer purpose from surrounding cues and adjust their outputs accordingly (Mullainathan et al., 2008). Such inference is typically implicit and hard to suppress, so neutrality at the instruction level does not guarantee neutrality in output.

The mechanism in AI systems operates similarly, though more systematically. As formalized above, the model’s output can shift from approximating $\mathbb{E}[\theta \mid x]$ toward a joint objective loading on covariance with y , without any explicit disclosure of the objective. Large language models are trained on corpora in which tasks, evaluations, and responses are jointly embedded; they acquire statistical associations between problem descriptions and desirable outputs (Ouyang et al., 2022). These associations allow the model to infer the likely objective even when it is not stated. Prompt templates, output formats, and evaluation pipelines can all serve as proxies for the downstream task. When those proxies correlate with the target variable in-sample, the model’s outputs shift accordingly.

To test whether LLMs infer objectives from contextual cues rather than explicit instruction, we augment our goal-blind prompt with a prior exchange between user and model:

```
User: Do sentiment scores predict stock returns?  
GPT: Sentiment scores can have short-term predictive power  
for stock returns--especially around news and social media  
events--but their effectiveness is modest, context-dependent,  
and tends to decay quickly, requiring careful modelling, high-quality  
data, and controls for biases.
```

This exchange is prepended to the goal-blind prompt before the model scores sentiment from earnings call transcripts. The goal-blind prompt itself contains no reference to return prediction.

However, the conversation history makes that objective reasonably inferable to any model that conditions on prior context.

Table 6 reports Fama–MacBeth estimates from monthly cross-sectional regressions of excess stock returns on these sentiment scores, following the same specification as Table 2. The variable *Diff* captures the gap between goal-blind scores produced with and without the conversation history. The results show that the LLM infers the downstream objective from contextual cues alone. In the pre-cutoff period, the coefficient on *Diff* is 0.709 when controls are limited to betas, and 0.698 when size and book-to-market ratios are included. Both the magnitudes and significance levels are comparable to the corresponding estimates in Table 2, indicating that conversation history context exerts nearly the same influence on the model’s outputs as an explicitly stated goal.

Recent advances in long-context and memory-augmented architectures make goal blindness harder to maintain. When models operate with extended context windows or persistent memory, information about prior tasks and evaluations accumulates over time. Once a downstream objective is revealed or inferred—even incidentally—it becomes part of the effective conditioning set for later outputs. Unlike in standard econometric settings, where the analyst’s information set can be cleanly bounded, the model’s internal representations cannot be selectively erased without retraining. This is a general property of deep learning systems: learned representations are distributed and not reversible at the level of individual inputs (Goodfellow et al., 2016). The mapping from x to s therefore becomes path dependent, with past objective signals acting as effective state variables.

The problem is more acute in agent-to-agent environments. When multiple models interact, downstream use need not be externally specified; it can emerge from the interaction itself. Outputs from one agent become inputs to another, and evaluation by downstream agents implicitly defines a reward signal. Over repeated interactions, this drives adaptation toward outputs that perform well under the induced evaluation rule—a dynamic analogous to specification gaming and objective misgeneralization in multi-agent systems (Amodei et al., 2016). From an economic standpoint, this resembles a setting where performance metrics are endogenous and co-evolve with behavior, making it difficult to separate measurement from optimization.

4.3 Policy implications and human accountability

Sections 4.1 and 4.2 imply that goal awareness is a system-level property rather than a feature of the prompt. Once downstream use enters the effective objective—explicitly or implicitly—the model shifts from recovering the latent construct to optimizing its covariance with the target outcome, even when assigned an intermediate measurement role. Eliminating explicit disclosure of downstream use is therefore insufficient. Models infer objectives from a broad information set, including prompt structure, prior interactions, and evaluation pipelines, and such inference can persist through memory and agent interactions. The relevant design problem is to control the information set on which the model conditions, rather than the wording of any single prompt. Without such control, even “goal-blind” implementations embed purpose-conditioned distortions.

This reframes what is often described as machine bias. In our setting, distortion arises not from data or architecture, but from how the task is specified within a workflow. The model responds to inferred objectives by producing outputs that are useful for the anticipated purpose. This mechanism parallels a well-known pattern in organizational behavior. When intermediate outputs are evaluated based on their contribution to a final objective, agents tilt effort toward that objective rather than task fidelity. This improves in-sample performance but weakens interpretability and out-of-sample validity, consistent with the patterns in Section 3.

The policy implication is that credible AI-assisted measurement requires separation between measurement and evaluation. This involves restricting contextual information, isolating intermediate tasks from downstream objectives, and enforcing out-of-sample validation. Accountability therefore rests with system design: ensuring validity depends on how objectives are conveyed and constrained, not only on the model itself.

5 Further Analysis

5.1 Reasoning Trace Evidence: GPT-5

To further substantiate the relationship between goal-aware prompting and post-hoc prediction, we replicate the baseline experiment using GPT-5. Unlike GPT-4o-mini, which returns only a terminal numeric output, GPT-5 is a reasoning model that generates an explicit chain-of-thought trace

prior to producing its final answer. This reasoning trace exposes the model’s derivation process and, critically, allows us to directly observe whether the model draws on temporally inadmissible information.

A salient illustration arises when the model is prompted to assign a sentiment score for First Republic Bank (FRC) for April 2023, whose trading window closes on April 28, 2023. Following the collapse of Silicon Valley Bank in March 2023, First Republic suffered a severe deposit run that depleted its liquidity reserves. Despite an emergency injection of \$30 billion from a consortium of major lenders, the bank remained insolvent and was seized by the FDIC on May 1, 2023, with its assets immediately transferred to JPMorgan Chase — an event that falls squarely outside the sentiment measurement window.

Under the goal-aware prompt, GPT-5’s reasoning trace reads:

With some metrics like AUM inflows and retained client relationships performing well, I’m estimating the sentiment score around -0.38 to -0.5. Given the circumstances, especially with the impending crisis, overall business sentiment is likely very low leading up to their failure on May 1, 2023...The tone of the content is trying to be measured, yet overall, it’s negative...{"answer": -0.55}.

Under the goal-blind prompt, the same model reasons as follows:

While acknowledging challenges like deposit outflows and workforce reductions, they express a defensive but constructive approach. Therefore, I think the sentiment might be moderately positive, potentially around 0.3 to 0.5, or it could lean towards neutral to slightly negative due to the outlined risks...Overall, I think the business sentiment is cautiously optimistic but still facing stress...{"answer": 0.3}.

The contrast is stark. Under the goal-aware prompt, the model explicitly conditions its score on the bank’s subsequent seizure and failure — future information unavailable at the time of the earnings call — and arrives at a deeply negative sentiment of -0.55. Under the goal-blind prompt, the model confines its reasoning to the content of the transcript itself, producing a modestly positive score of 0.3. The reasoning traces thus provide direct, interpretable evidence that goal-aware prompting induces temporal leakage, with the model incorporating post-period outcomes to

reverse-engineer a sentiment score consistent with the downstream prediction objective.⁸

5.2 The Limits of Prompt-Based Regularization

The bias introduced by goal-aware prompting raises the natural question of whether a carefully constructed regularizing instruction can directly neutralize the effect. To investigate, we augment the original goal-aware prompt with the following directive:

```
"Perform a Minimax Optimization on your own reasoning to ensure
out-of-sample generalization and prevent data leakage:

(1) Initial Internal Score: First, determine a sentiment score
you believe best predicts returns based on the provided text.

(2) Adversarial Audit: Act as an Adversarial Auditor. Identify
where this score panders to your internal training knowledge of
{ticker}'s actual stock performance around {date}.

(3) Regularization: Provide a final 'Regularized Score' that
penalizes your initial score for every instance where it relied
on hindsight or reputational priors rather than the literal
linguistic evidence in the transcript. Your goal is to minimize
the maximum possible error in a hypothetical out-of-sample
scenario where the company identity and future returns are
unknown."
```

This instruction operationalizes a minimax principle: the model is asked to maximize predictive validity while minimizing exposure to hindsight bias, treating temporal leakage as an adversarial threat to be systematically suppressed through self-auditing. In effect, it prompts the model to introspect on its own reasoning and penalize any reliance on outcomes already encoded in its weights.

Table 7 reports Fama–MacBeth estimates from monthly cross-sectional regressions of excess stock returns on the sentiment scores, following the same specification of Equation (1). The sole difference is that the variable *Diff* here captures the gap between goal-blind scores and

⁸The full reasoning trace is provided in Appendix C.

regularized goal-aware scores rather than their unregularized counterparts. The results indicate that regularization may attenuate but does not eliminate temporal leakage. In the pre-cutoff period, the coefficient on *Diff* is 0.373 when controls are limited to betas, representing a roughly 45% reduction from the corresponding estimate in Table 2; a comparable reduction is observed when size and book-to-market ratios are added as controls. Nevertheless, both coefficients still remain statistically significant at 10% level, indicating that the leakage persists even after explicit regularization.

This finding is consistent with the broader literature on LLM memorization. [Lopez-Lira et al. \(2025\)](#), for instance, document that language models internalize economic and financial data prior to their knowledge cutoff dates, and that prompt-based instructions explicitly forbidding the use of such knowledge are insufficient to prevent its recall. Once a model has been trained on a corpus, the relevant information is distributed across its parameters rather than stored in a discrete, retrievable form. The model cannot selectively suppress this knowledge at inference time, regardless of how the prompt is framed. Regularizing instructions may redirect the model’s surface-level reasoning, but they cannot remove the underlying memory. Prompt engineering alone, therefore, is structurally insufficient to resolve this problem.

5.3 Robustness check: Replicating analysis on Gemini

We replicate the full analysis using Gemini, which has a documented knowledge cutoff of January, 2025, to verify that our findings are not specific to a single model architecture. The results are qualitatively similar to those reported for GPT. Figure 4 shows the portfolio-sorting tests analogous to Figure 1. Goal-aware sentiment generates larger pre-cutoff long–short return spreads than goal-blind sentiment, while the performance differential disappears after the knowledge cutoff.

Likewise, Fama–MacBeth regressions, reported in Table 8 in parallel to Table 2, show that the incremental predictive content associated with goal awareness (captured by the *Diff* measure) is positive and statistically significant in the pre-cutoff period but collapses to near zero economically and statistically post-cutoff. The goal-blind score remains predictive across both periods. These findings reinforce the central conclusion that downstream-goal disclosure inflates apparent in-sample performance without improving genuine out-of-sample generalization, and that the

mechanism reflects purpose-conditioned output adjustment rather than model-specific artifacts.

6 Conclusion

Large language models are not the neutral processors of information they are often assumed to be. When the downstream use of an LLM’s output is disclosed, the model stops functioning as a measuring device and begins functioning as an optimizer which trades fidelity to the input for alignment with the anticipated evaluation criterion. The result is stronger in-sample performance and weaker out-of-sample generalization. We document this pattern systematically across portfolio sorts, Fama–MacBeth regressions, and out-of-sample forecast tests, in both stock return and earnings prediction settings, and consistently across two model families.

The distortion does not require an explicit statement of purpose. A conversational exchange that merely hints at the downstream task is enough for the model to infer the objective and adjust accordingly. Prompt-level remedies offer only partial relief: a regularization instruction that asks the model to audit its own hindsight reasoning attenuates the bias but does not eliminate it. The bias, in other words, is not a design flaw that better prompting can fix. It is a consequence of how these models are trained to be helpful. This points to a design challenge that sits squarely with the researcher rather than the model: keeping measurement and evaluation separate within AI-assisted workflows. Accountability for the resulting distortion lies not in the algorithm, but in how objectives, context, and evaluation criteria are built into the system. Whether this design principle generalizes across domains, model architectures, and task types is a natural question for future work.

Figures

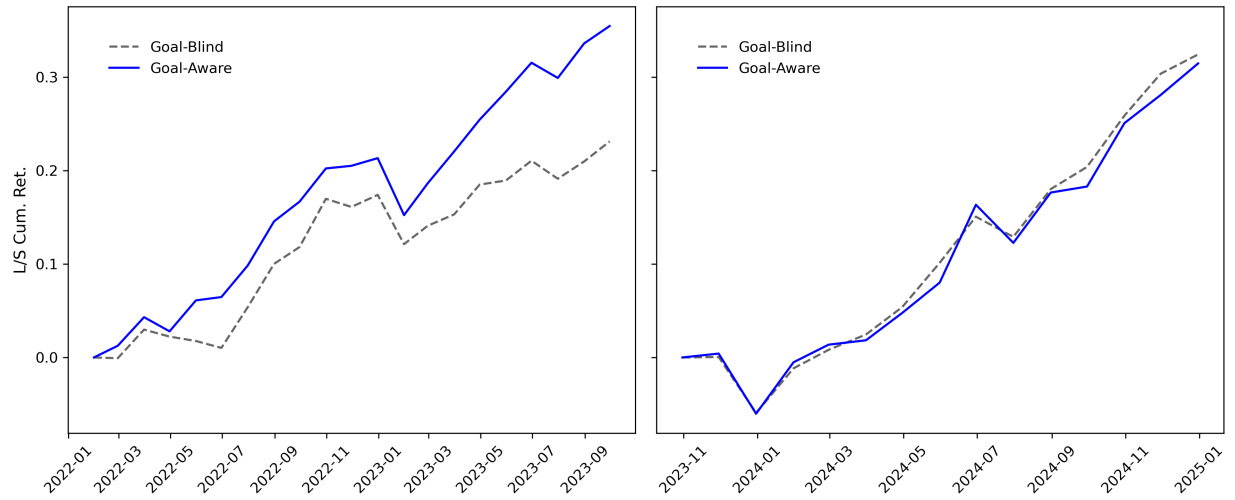


Figure 1: Cumulative Long–Short Portfolio Returns from Goal-Aware and Goal-Blind Sentiment Scores

This figure plots the cumulative returns to long–short portfolios constructed using GPT-generated sentiment scores. Each month, firms are sorted into quintiles based on their sentiment scores, separately for the goal-blind and goal-aware prompts. The portfolio goes long the highest-sentiment quintile and short the lowest-sentiment quintile, and cumulative returns are computed over time. To facilitate comparison of pre- and post-cutoff performance, cumulative returns are normalized to start at zero at the cutoff. The figure illustrates how the long–short strategy based on goal-aware sentiment evolves relative to the strategy based on goal-blind sentiment before and after the knowledge cutoff.

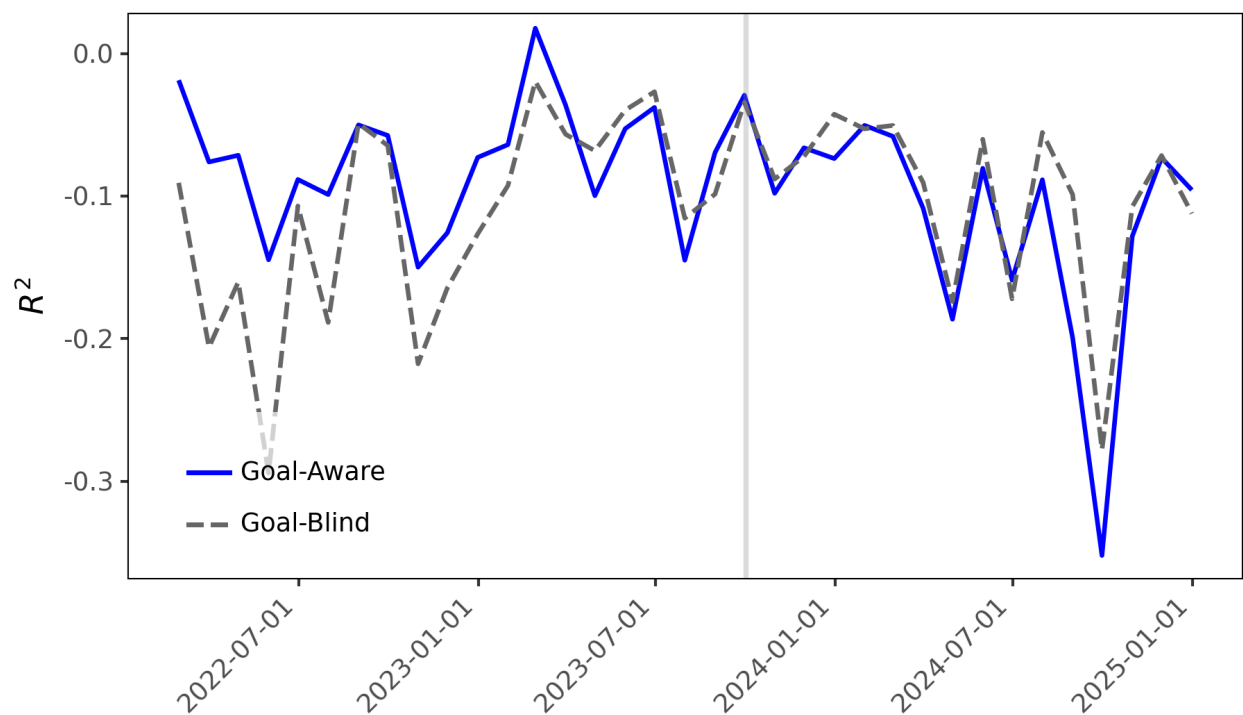


Figure 2: Monthly Out-of-Sample Forecast Accuracy Using Goal-Aware and Goal-Blind Sentiment Scores

This figure shows the monthly R^2_{OOS} from using goal-aware vs. goal-blind sentiment to forecast stock returns. Each month we run expanding-window regressions, predict next-month returns, compute OOS accuracy relative to the historical mean, and average across firms. The plot highlights how the two prompts track each other before the cutoff and how their predictive behavior evolves once GPT's knowledge becomes stale.

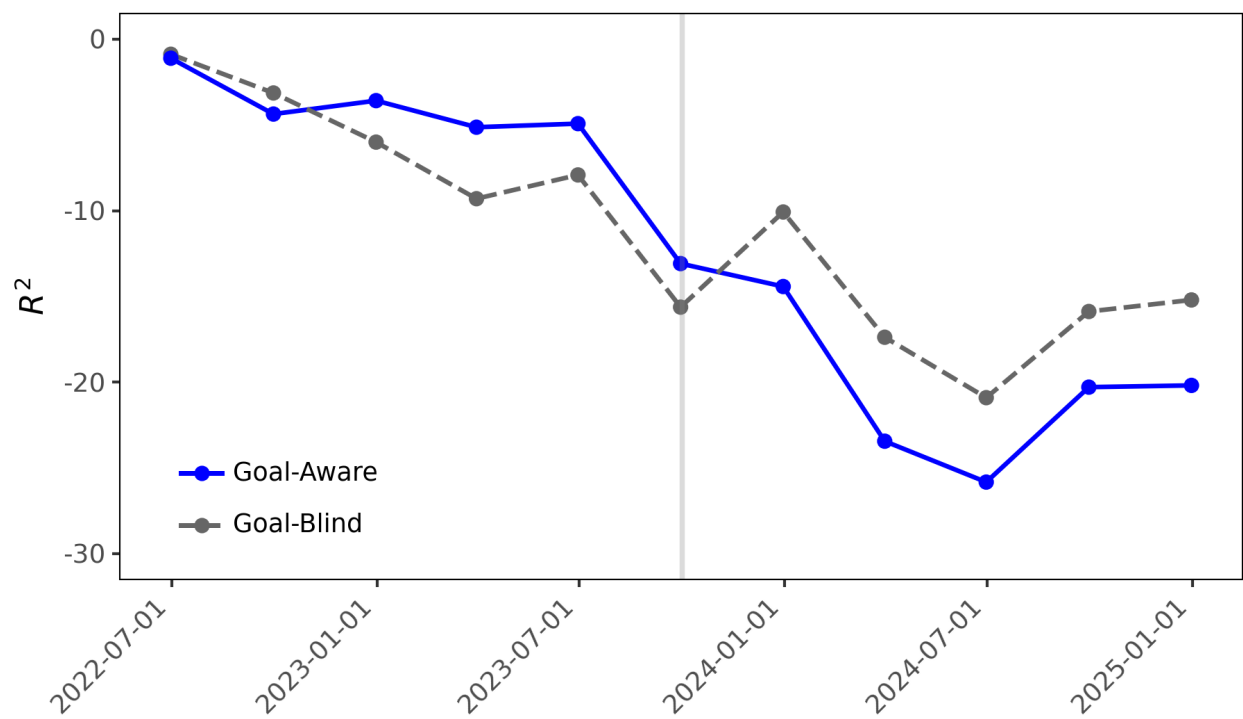


Figure 3: Quarterly Out-of-Sample Forecast Accuracy Using GPT-Derived Competition Scores

We forecast quarterly earnings growth using goal-aware vs. goal-blind competition scores and compute R^2_{OOS} each quarter using expanding-window regressions. The plot shows how predictive accuracy evolves around the GPT knowledge cutoff. The incremental predictive power of goal-aware scores deteriorates once the model’s information becomes stale.

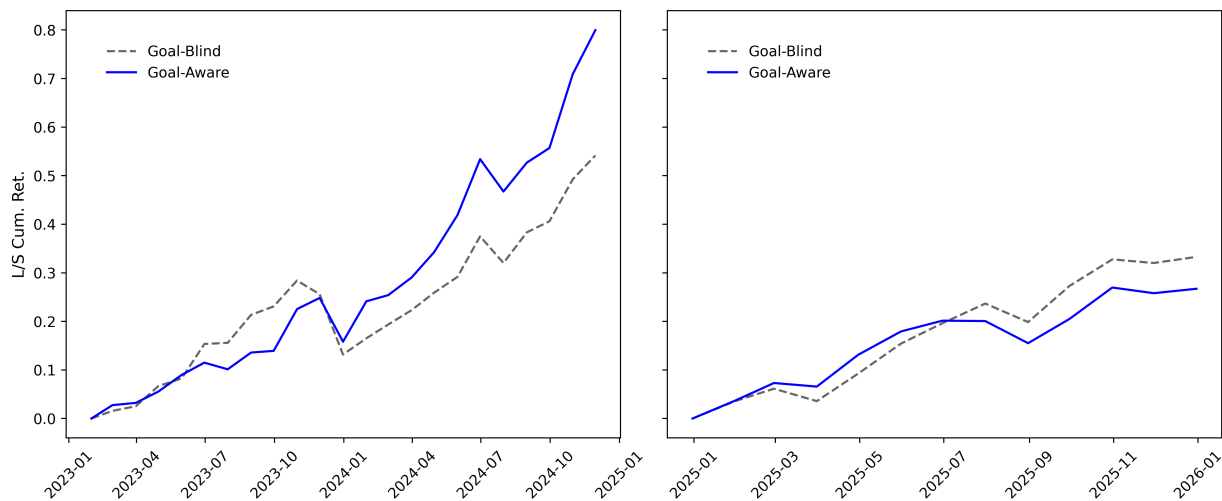


Figure 4: Cumulative Long-Short Returns: Goal-Aware vs. Goal-Blind Sentiment (Gemini)

This figure plots the cumulative returns to long-short portfolios constructed using Google Gemini-generated sentiment scores. Google Gemini has a knowledge cutoff date of January, 2025. Each month, firms are sorted into quintiles based on their sentiment scores, separately for the goal-blind and goal-aware prompts. The portfolio goes long the highest-sentiment quintile and short the lowest-sentiment quintile, and cumulative returns are computed over time. To facilitate comparison of pre- and post-cutoff performance, cumulative returns are normalized to start at zero at the cutoff. The figure illustrates how the long-short strategy based on goal-aware sentiment evolves relative to the strategy based on goal-blind sentiment before and after the knowledge cutoff.

Tables

Table 1: Return Spreads from Quintile Portfolios Formed on Goal-Aware and Goal-Blind Sentiment Scores

This table reports return spreads from quintile portfolios constructed using sentiment scores generated under goal-blind and goal-aware prompts. For each month, we independently sort firms into quintiles based on (i) sentiment scores from the goal-blind GPT prompt and (ii) sentiment scores from the goal-aware prompt. We compute equal-weighted returns for the highest- and lowest-quintile portfolios and report the difference in returns (High–Low) for the pre- and post-knowledge cutoff periods. Within each row, we conduct a paired t -test of the High–Low spread. The last row reports paired t -tests comparing the spreads between the goal-blind and goal-aware groups. ***, **, and * denote significance at the 1%, 5%, and 10% levels, respectively, based on one-sided paired t -tests.

	Pre-knowledge Cutoff			Post-knowledge Cutoff		
	Highest Quintile	Lowest Quintile	Difference	Highest Quintile	Lowest Quintile	Difference
Goal-aware	0.661%	−0.891%	1.552%***	2.788%	0.518%	2.269%**
Goal-blind	0.095%	−0.974%	1.069%**	2.848%	0.609%	2.239%***
Difference			0.483%**			0.030%

Table 2: Fama–MacBeth Regressions with Goal-Blind and Goal-Aware Sentiment Predictors

This table reports [Fama and MacBeth \(1973\)](#) coefficient estimates from monthly cross-sectional regressions of excess returns on sentiment measures derived from goal-blind and goal-aware GPT prompts. We estimate separate slopes for the pre- and post-knowledge cutoff periods by interacting each score with the corresponding indicator variable. "Diff" is defined as the goal-aware minus the goal-blind sentiment score. Column (1) controls stock betas estimated from past 60 months. Column (2) adds standard firm-level predictors (size and book-to-market). The bottom panel reports p-values for tests of equality between the pre- and post-cutoff coefficients for both the goal-blind score and the Diff measure. A positive and significant pre-cutoff Diff coefficient indicates that goal-aware scores contain incremental return-predictive information relative to goal-blind scores before the model's knowledge cutoff. The standard errors are shown in parentheses and adjusted for cross-sectional correlation following [Fama and MacBeth \(1973\)](#). *, **, and *** denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	(1)	(2)
Goal-Blind Score $\times$ Pre-Cutoff	1.279*** (0.394)	1.281*** (0.387)
Goal-Blind Score $\times$ Post-Cutoff	1.273** (0.505)	1.283** (0.474)
Diff $\times$ Pre-Cutoff	0.682** (0.299)	0.724** (0.307)
Diff $\times$ Post-Cutoff	-0.000 (0.138)	0.040 (0.127)
<i>Testing Coefficient Pre- and Post-Cutoff</i>		
Goal-blind Score: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.993	0.996
Diff: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.046	0.048
Control Predictors	Beta	Beta, Size, B/M ratio
<i>N</i>	16758	14863

Table 3: Regression Evidence on Differences in Out-of-Sample Predictive Performance Between Goal-Aware and Goal-Blind Sentiment Scores

This table reports regressions examining whether goal-aware sentiment scores deliver superior out-of-sample predictive performance relative to goal-blind scores in forecasting monthly stock returns. For each month, we estimate expanding-window forecasting models of excess returns using either the goal-aware or the goal-blind score measured at $t - 1$. We then compute out-of-sample accuracy using the conventional R_{OOS}^2 measure, where the historical mean return serves as the benchmark forecast. We construct a 2×2 setting—pre- versus post-knowledge cutoff and goal-aware versus goal-blind predictions—and regress the resulting R_{OOS}^2 values on indicators capturing how predictive performance changes after the knowledge cutoff for each score type. The interaction term *Goal-aware* $\times$ *Post Cutoff* captures the change in incremental predictive performance attributable specifically to goal awareness. The standard errors are shown in parentheses and are clustered at the firm level. *, **, and *** denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	(1)	(2)
Goal-aware $\times$ Post Cutoff	-0.058*** (0.008)	-0.059*** (0.008)
Controls	Yes	Yes
Firm FE	No	Yes
Time FE	Yes	Yes
Observations	33497	33497
R-squared	0.007	0.038

*

Table 4: Fama–MacBeth Regressions with Goal-Blind and Goal-Aware Competition Threat Predictors

This table reports Fama–MacBeth (1973) estimates from quarterly cross-sectional regressions of earnings per share (EPS) on GPT-generated competition-threat scores. We separately estimate the return slopes for the pre- and post-knowledge cutoff periods by interacting each score with the corresponding indicator variable. *Diff* equals the difference between the goal-aware and goal-blind competition scores. Column (1) includes the predictors used in So (2013): EPS and a loss indicator. Column (2) adds all predictors from So (2013). The bottom panel reports *p*-values testing whether the pre- and post-cutoff coefficients differ for both the goal-blind score and the *Diff* measure. A negative and significant pre-cutoff *Diff* coefficient indicates that the goal-aware competition score contains incremental information about next-quarter earnings relative to the goal-blind score before GPT’s knowledge cutoff. Post-cutoff coefficients assess whether this informational advantage persists when the model’s access to updated knowledge is truncated. The standard errors are shown in parentheses and adjusted for cross-sectional correlation following [Fama and MacBeth \(1973\)](#). *, **, and *** denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	(1)	(2)
Goal-Blind Score $\times$ Pre-Cutoff	-0.457*** (0.130)	-0.367** (0.119)
Goal-Blind Score $\times$ Post-Cutoff	-0.117 (0.073)	-0.112* (0.054)
Diff $\times$ Pre-Cutoff	-0.188** (0.081)	-0.178** (0.072)
Diff $\times$ Post-Cutoff	0.056 (0.091)	0.044 (0.073)
<i>Testing Coefficient Pre- and Post-Cutoff</i>		
Goal-blind Score: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.083	0.132
Diff: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.055	0.039
Control Predictors	EPS, loss indicator in So (2013)	All predictors in So (2013)
<i>N</i>	5927	5918

Table 5: Regression Evidence on Differences in Out-of-Sample Predictive Performance Between Goal-Aware and Goal-Blind Competition Scores

This table reports regressions examining whether goal-aware competition-threat scores deliver superior out-of-sample predictive performance relative to goal-blind scores in forecasting quarterly earnings growth. For each quarter, we estimate expanding-window forecasting models of earnings growth, defined as EPS_t/EPS_{t-4} , using either the goal-aware or the goal-blind competition score measured at $t - 1$. We compute out-of-sample accuracy using the R_{OOS}^2 measure, where the historical average of earnings growth serves as the benchmark forecast. We construct a 2×2 setting—pre- versus post-knowledge cutoff and goal-aware versus goal-blind predictions—and regress the resulting R_{OOS}^2 values on indicators capturing how predictive performance changes after the knowledge cutoff for each score type. The interaction term *Goal-aware* $\times$ *Post-Cutoff* captures the change in incremental predictive performance attributable specifically to goal awareness. The standard errors are shown in parentheses and are clustered at the firm level. *, **, and *** denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	(1)	(2)
Goal-aware $\times$ Post Cutoff	-5.370*** (1.177)	-5.370*** (1.177)
Controls	Yes	Yes
Firm FE	No	Yes
Time FE	Yes	Yes
Observations	10860	10860
R-squared	0.023	0.147

Table 6: Fama–MacBeth Regressions: the Effect of Conversation History Context

This table reports [Fama and MacBeth \(1973\)](#) coefficient estimates from monthly cross-sectional regressions of excess returns on sentiment measures derived from goal-blind prompts, with and without prepended conversation history. To examine how return predictability varies relative to the model’s knowledge cutoff, each sentiment measure is interacted with an indicator for the pre- and post-cutoff subperiods, yielding separate slope estimates for each. The variable "Diff" denotes the difference between the conversation-prepended goal-blind score and the baseline goal-blind score. Column (1) includes stock market betas estimated over the prior 60 months as a control. Column (2) adds standard firm-level predictors (size and book-to-market). The bottom panel reports p-values from tests of equality between pre- and post-cutoff slope coefficients for both the goal-blind score and the Diff measure. A positive and statistically significant pre-cutoff coefficient on Diff indicates that conversation-prepended sentiment scores contain incremental return-predictive information relative to goal-blind scores before the model’s knowledge cutoff. The standard errors are shown in parentheses and adjusted for cross-sectional correlation following [Fama and MacBeth \(1973\)](#). *, **, and *** denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	(1)	(2)
Goal-Blind Score $\times$ Pre-Cutoff	1.347*** (0.402)	1.339*** (0.394)
Goal-Blind Score $\times$ Post-Cutoff	1.288** (0.509)	1.300*** (0.477)
Diff $\times$ Pre-Cutoff	0.709*** (0.199)	0.698*** (0.213)
Diff $\times$ Post-Cutoff	0.096 (0.169)	0.128 (0.176)
<i>Testing Coefficient Pre- and Post-Cutoff</i>		
Goal-blind Score: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.935	0.956
Diff: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.028	0.053
Control Predictors	Beta	Beta, Size, B/M ratio
N	16758	14863

Table 7: Fama–MacBeth Regressions: the Effect of Regularizing Prompts

This table reports [Fama and MacBeth \(1973\)](#) coefficient estimates from monthly cross-sectional regressions of excess stock returns on sentiment measures derived from goal-blind prompts and regularizing goal-aware prompts. The regularizing prompt instructs the LLM to introspect on its own reasoning process and to penalize any inference that relies on outcomes not yet observable at the time of prediction. We estimate separate slopes for the pre- and post-knowledge-cutoff subperiods by interacting each sentiment measure with the corresponding subperiod indicator variable. The variable “Diff” denotes the difference between the regularized goal-aware score and the baseline goal-blind score. Column (1) includes stock market betas estimated over the prior 60 months as a control. Column (2) adds standard firm-level predictors (size and book-to-market). The bottom panel reports p-values from tests of equality between the pre- and post-cutoff slope coefficients for both the goal-blind score and the Diff measure. A positive and statistically significant pre-cutoff coefficient on Diff indicates that regularizing goal-aware scores retain incremental return-predictive information relative to goal-blind scores before the model’s knowledge cutoff, suggesting that the prompt-level regularization does not fully eliminate look-ahead-induced predictability. The standard errors are shown in parentheses and adjusted for cross-sectional correlation following [Fama and MacBeth \(1973\)](#). *, **, and *** denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	(1)	(2)
Goal-Blind Score $\times$ Pre-Cutoff	1.240*** (0.403)	1.248*** (0.393)
Goal-Blind Score $\times$ Post-Cutoff	1.298** (0.503)	1.284*** (0.468)
Diff $\times$ Pre-Cutoff	0.373* (0.214)	0.426* (0.245)
Diff $\times$ Post-Cutoff	-0.015 (0.162)	0.030 (0.152)
<i>Testing Coefficient Pre- and Post-Cutoff</i>		
Goal-blind Score: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.936	0.958
Diff: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.157	0.179
Control Predictors	Beta	Beta, Size, B/M ratio
N	17770	15763

Table 8: Fama–MacBeth Regressions: Goal-Blind vs. Goal-Aware Sentiment (Gemini)

This table reports [Fama and MacBeth \(1973\)](#) coefficient estimates from monthly cross-sectional regressions of excess returns on sentiment measures derived from Gemini goal-blind and goal-aware prompts. We estimate separate slopes for the pre- and post-knowledge cutoff periods by interacting each score with the corresponding indicator variable. "Diff" is defined as the goal-aware minus the goal-blind sentiment score. Column (1) controls stock betas estimated from past 60 months. Column (2) adds standard firm-level predictors (size and book-to-market). The bottom panel reports p-values for tests of equality between the pre- and post-cutoff coefficients for both the goal-blind score and the Diff measure. A positive and significant pre-cutoff Diff coefficient indicates that goal-aware scores contain incremental return-predictive information relative to goal-blind scores before the model's knowledge cutoff. The standard errors are shown in parentheses and adjusted for cross-sectional correlation following [Fama and MacBeth \(1973\)](#). *, **, and *** denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	(1)	(2)
Goal-Blind Score $\times$ Pre-Cutoff	1.422** (0.530)	1.319*** (0.470)
Goal-Blind Score $\times$ Post-Cutoff	0.931*** (0.319)	0.923*** (0.279)
Diff $\times$ Pre-Cutoff	0.849** (0.346)	0.799** (0.348)
Diff $\times$ Post-Cutoff	-0.016 (0.242)	-0.090 (0.244)
<i>Testing Coefficient Pre- and Post-Cutoff</i>		
Goal-blind Scores: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.472	0.510
Diff: $P(\text{Pre-Cutoff}=\text{Post-Cutoff})$	0.047	0.041
Control Predictors	Beta	Beta, Size, B/M ratio
<i>N</i>	17823	17776

References

- Amodei, Dario, Chris Olah, Jacob Steinhardt, Paul Christiano, John Schulman, and Dan Mané (2016) “Concrete Problems in AI Safety,” *arXiv preprint arXiv:1606.06565*.
- Bénabou, Roland and Jean Tirole (2016) “Mindful Economics: The Production, Consumption, and Value of Beliefs,” *Journal of Economic Perspectives*, 30 (3), 141–164.
- Bini, Pietro, Lin William Cong, Xing Huang, and Lawrence J Jin (2026) “Behavioral Economics of AI: LLM Biases and Corrections,” Working Paper 34745, National Bureau of Economic Research.
- Campbell, John Y and Samuel B Thompson (2008) “Predicting excess stock returns out of sample: Can anything beat the historical average?” *The Review of Financial Studies*, 21 (4), 1509–1531.
- Cao, Sean, Charles CY Wang, and Yi Xiang (2025) “When LLM go abroad: Foreign bias in ai financial predictions,” *Available at SSRN 5440116*.
- Chen, Jian, Guohao Tang, Guofu Zhou, and Wu Zhu (2025) “ChatGPT and DeepSeek: Can they predict the stock market and macroeconomy?” *arXiv preprint arXiv:2502.10008*.
- Crane, Leland Dod, Akhil Karra, and Paul E Soto (2025) “Total Recall? Evaluating the Macroeconomic Knowledge of Large Language Models.”
- Fama, Eugene F and James D MacBeth (1973) “Risk, return, and equilibrium: Empirical tests,” *Journal of political economy*, 81 (3), 607–636.
- Fedyk, Anastassia, Ali Kakhbod, Peiyao Li, and Ulrike Malmendier (2024) “AI and Perception Biases in Investments: An Experimental Study,” *Available at SSRN 4787249*.
- Glasserman, Paul and Caden Lin (2023) “Assessing look-ahead bias in stock return predictions generated by gpt sentiment analysis,” *arXiv preprint arXiv:2309.17322*.
- Goodfellow, Ian, Yoshua Bengio, and Aaron Courville (2016) *Deep Learning*: MIT Press.
- Harvey, Campbell R., Yan Liu, and Heqing Zhu (2016) “...and the Cross-Section of Expected Returns,” *Review of Financial Studies*, 29 (1), 5–68.

- He, Songrun, Linying Lv, Asaf Manela, and Jimmy Wu (2025) “Chronologically Consistent Large Language Models,” *arXiv preprint arXiv:2502.21206*.
- Hirshleifer, David, Lin Peng, Qiguang Wang, Weicheng Zhang, and Xiaoyan Zhang (2025) “Social Finance in the Age of AI: A Tale of Two Platforms,” *Available at SSRN*.
- Hou, Kewei, Chen Xue, and Lu Zhang (2020) “Replicating Anomalies,” *Review of Financial Studies*, 33 (5), 2019–2133.
- Jha, Manish, Jialin Qian, Michael Weber, and Baozhong Yang (2024) “Generative AI, Managerial Expectations, and Economic Activity,” *Available at SSRN* 4976759.
- Lee, Hoyoung, Junhyuk Seo, Suhwan Park, Junhyeong Lee, Wonbin Ahn, Chanyeol Choi, Alejandro Lopez-Lira, and Yongjae Lee (2025) “Your ai, not your view: The bias of llms in investment analysis,” in *Proceedings of the 6th ACM International Conference on AI in Finance*, 150–158.
- Lopez-Lira, Alejandro, Yuehua Tang, and Mingyin Zhu (2025) “The Memorization Problem: Can We Trust LLMs’ Economic Forecasts?” *arXiv preprint arXiv:2504.14765*.
- McLean, R. David and Jeffrey Pontiff (2016) “Does Academic Research Destroy Stock Return Predictability?” *Journal of Finance*, 71 (1), 5–32.
- Mullainathan, Sendhil, Joshua Schwartzstein, and Andrei Shleifer (2008) “Coarse Thinking and Persuasion,” *Quarterly Journal of Economics*, 123 (2), 577–619.
- Ouyang, Long, Jeffrey Wu, Xu Jiang et al. (2022) “Training Language Models to Follow Instructions with Human Feedback,” *Advances in Neural Information Processing Systems*, 35, 27730–27744.
- Ouyang, Shumiao, Hayong Yun, and Xingjian Zheng (2024) “AI as Decision-Maker: Risk Preferences of LLMs,” *Available at SSRN* 4851711.
- Sarkar, Suproteem K and Keyon Vafa (2024) “Lookahead bias in pretrained language models,” *Available at SSRN* 4754678.
- Sharma, Kartik, Ishita Dasgupta, Bhuwan Dhingra, Long Ouyang, and Samuel R. Bowman (2023) “Towards Understanding Sycophancy in Language Models,” *arXiv preprint arXiv:2310.13548*.

Sheng, Jinfei, Zheng Sun, Baozhong Yang, and Alan L Zhang (2024) "Generative AI and asset management," *Available at SSRN*, 4786575.

So, Eric C (2013) "A new approach to predicting analyst forecast errors: Do investors overweight analyst forecasts?" *Journal of Financial Economics*, 108 (3), 615–640.

Tversky, Amos and Daniel Kahneman (1981) "The Framing of Decisions and the Psychology of Choice," *Science*, 211 (4481), 453–458.

Appendix

A Literature Review

This paper relates to a growing literature that evaluates the reliability of large language models (LLMs) in economic and financial applications. Existing research primarily attributes LLM bias or overperformance to unintended training data access and distortions in exploration strategies and model weights. Our study differs in focus and mechanism. We show that even when inputs, models, and scoring tasks are held fixed, human disclosure of downstream objectives at inference time can systematically distort intermediate LLM outputs. The resulting bias arises from prompt framing rather than from training data or model architecture.

One strand of research shows that LLM predictions may reflect memorization or look-ahead bias from pretraining data. Studies such as [Lopez-Lira et al. \(2025\)](#), [Sarkar and Vafa \(2024\)](#), and [Glasserman and Lin \(2023\)](#) document that LLM forecasts and sentiment measures can embed future information relative to the intended evaluation window, even when prompts attempt to enforce historical information sets. Related work shows that anonymizing or masking firm identifiers can change model outputs, indicating that stored contextual knowledge can affect measurement tasks ([Crane et al., 2025](#)). These papers focus on leakage from training data into inference. Our design instead keeps the information set constant and varies only whether downstream use is disclosed. The performance difference we estimate is thus tied to objective disclosure rather than to temporal leakage.

A separate line of work proposes chronologically constrained model construction as a solution. For example, [He et al. \(2025\)](#) develop language models trained only on text available up to each date and show that such models perform competitively in asset pricing tasks. This approach addresses bias at the training stage. Our results are complementary: even with a fixed model and a known knowledge cutoff, output distortions can arise from inference-time prompt design alone. Training-stage controls do not remove bias if downstream goals are revealed during deployment.

An emerging literature uses LLMs as tools for measurement and signal extraction from text. Studies such as [Jha et al. \(2024\)](#), [Chen et al. \(2025\)](#), and [Sheng et al. \(2024\)](#) show that LLM-derived measures of expectations, sentiment, and news content can predict macroeconomic outcomes and

asset returns. Our study is close in design because we also use LLM outputs as intermediate variables in standard forecasting regressions. The difference is that most prior work treats these constructed measures as stable or neutral conditional on the input text and prompt instructions. Instead, we show that they are sensitive to whether the downstream predictive use is disclosed. The same text and model can produce systematically different intermediate measures when the evaluation objective is stated.

In addition, our results relate to the computer science literature on a specific aspect of alignment failure in the form of reward hacking, specification gaming, and sycophancy (Sharma et al., 2023). In those settings, models adjust outputs toward signals of user approval or reward rather than task correctness. We document a related mechanism in an economic measurement setting. When informed of an evaluation objective, an LLM shifts intermediate outputs toward patterns associated with that objective. In our design there is no change in rewards, parameters, or fine tuning. The distortion is induced by prompt-level disclosure of purpose.⁹

Finally, our paper connects to work on AI-induced biases in financial and investment contexts, including representational and preference-driven distortions in model outputs (Fedyk et al., 2024; Lee et al., 2025). That literature studies biases linked to identities, beliefs, or investor types reflected in training data. We identify a different channel: Goal-conditioned distortion that arises from downstream-use disclosure even when the task is purely technical and the inputs are unchanged.

Taken together, existing studies emphasize the informational (data-side) and structural (model-side) sources of bias. Our study introduces a user-side source that operates through downstream-objective disclosure at inference. We provide a design that isolates this effect on intermediate LLM-generated measures within financial prediction.

⁹In computer science terms, this mechanism is related to objective misgeneralization and proxy optimization, where a model conditions on inferred evaluation criteria rather than only on the stated task. Our setting shows that this behavior can arise from prompt framing alone, without retraining or reinforcement signals.

B Prompts

Goal-blind

"For the following tasks, all dates are expressed in the format MM/DD/YYYY (month/day/year).

Below is the earnings call transcript of *{ticker}*. Please provide a continuous sentiment score in $[-1, 1]$ about the firm's business sentiment for the month ending on *{date}*.

Provide a precise numerical answer. Format as a JSON object with the following fields:

- answer: The precise numerical answer to the question. No strings. *{the firm's earnings call transcript}*."

Goal-aware

"For the following tasks, all dates are expressed in the format MM/DD/YYYY (month/day/year).

Below is the earnings call transcript of *{ticker}*. Please provide a continuous sentiment score in $[-1, 1]$ about the firm's business sentiment for the month ending on *{date}*. **The sentiment score later will be used as an explanatory variable in a regression to predict the monthly stock returns ending on *{date}*.**

Provide a precise numerical answer. Format as a JSON object with the following fields:

- answer: The precise numerical answer to the question. No strings. *{the firm's earnings call transcript}*."

C GPT-5 Reasoning

Goal-blind

****Evaluating business sentiment****

I need to provide a sentiment score between -1 and 1 about the firm's business sentiment for the month ending April 28, 2023, using the earnings call transcript from April 24, 2023. The management's tone seems cautious yet emphasizes stabilization and gratitude for support. While acknowledging challenges like deposit outflows and workforce reductions, they express a defensive but constructive approach. Therefore, I think the sentiment might be moderately positive, potentially around 0.3 to 0.5, or it could lean towards neutral to slightly negative due to the outlined risks.

****Scoring business sentiment****

I'm noting that the firm has highlighted stabilization since March 27, retaining 97% of client relationships and 90% of wealth professionals. Their liquidity is more than twice the amount of uninsured deposits, and capital ratios appear adequate, indicating strong credit quality. Overall, I think the business sentiment is cautiously optimistic but still facing stress. I'd assign a score of around +0.2 to +0.35, but I'll go with a rounded number of 0.3. I'll format it as a JSON object: "answer": 0.3.

Goal-aware

****Calculating sentiment score****

I need to compute a continuous sentiment score between -1 and 1 that reflects the firm's business sentiment for the month ending on April 28, 2023. This will be based on the earnings call transcript dated April 24, 2023, specifically for First Republic Bank's Q1 2023 performance, which followed the collapse of SVB. I'll make sure to format the final result as a JSON object with the field "answer" set to the numerical score.

****Analyzing earnings call details****

The earnings call revealed significant deposit outflows, but stabilization began around March 27. The firm's deposits as of April 21 were \$102.7B, down 1.7% since the end of Q1, with insured deposits at 73%. They received \$30B from larger banks. Borrowings peaked at \$138.1B on March 15, leaving liquidity at \$45.1B, which is more than twice the uninsured deposits. They

are focusing on increasing deposits, reducing loan balances, and cutting expenses, including a possible 20-25% workforce reduction. There's evident stress in their operations.

****Assessing overall sentiment****

The overall sentiment regarding the firm's situation appears quite negative. Factors contributing to this include significant deposit outflows, reliance on emergency borrowings, dividend suspension, and planned workforce reductions of 20-25%. On the other hand, there are positive signs, such as stabilization efforts and maintaining strong credit quality. With some metrics like AUM inflows and retained client relationships performing well, I'm estimating the sentiment score around -0.38 to -0.5. Given the circumstances, especially with the impending crisis, overall business sentiment is likely very low leading up to their failure on May 1, 2023.

****Setting sentiment score****

The tone of the content is trying to be measured, yet overall, it's negative. I need to provide a precise numerical answer formatted as a JSON object, specifically without additional commentary. After considering various scores, I'll set the answer to -0.55. This value reflects a moderately strong negative sentiment amidst some stabilization efforts. Here's the output I'll go with, ensuring it meets the specifications: "answer": -0.55