

The AI Proxy Advisor¹

Michiel Bakker*, Maxime Couvert[†], Sercan Demir[‡], Roni Michaely[§]

May 2026

Abstract

We examine whether artificial intelligence (AI) can help shareholders make value-enhancing voting decisions. Using a causal machine learning framework that combines regression discontinuity with ensemble prediction models, we show that votes aligned with AI recommendations generate higher cumulative abnormal returns and lead to superior long-term firm performance. Proxy advisors not only do not appear to make value-enhancing recommendations but also fail to improve AI guidance. One reason shareholders fail to vote optimally is their excessive support for management proposals. By capturing complex nonlinear relationships between proposals and firm characteristics, AI consistently identifies strategies overlooked by both shareholders and proxy advisors.

JEL Classification: G20, G23, G30, G34, M14, Q56.

Keywords: Corporate governance, corporate social responsibility, proxy voting, proxy advisors, causal machine learning

¹We appreciate comments from Wei Jiang and conference participants at the HKU 2025 Sustainability Forum, HKU-Belgium 2025 Sustainability Conference, HKU-Oxford AI for Business Conference, 4th HKU Summer Finance Conference. We also thank the Cooperative AI Foundation for supporting this work. We thank Jonathon Zytnick for providing data on proxy advisors' recommendations.

* MIT Sloan School of Management and MIT Institute for Data, Systems, and Society: bakker@mit.edu

[†] HKU Business School at University of Hong Kong: mcouvert@hku.hk

[‡] University of Cologne: demir@wiso.uni-koeln.de

[§] HKU Business School at University of Hong Kong: ronim@hku.hk

1 Introduction

Institutional investors routinely vote on a wide range of proposals at their portfolio firms, including director elections, executive compensation, and environmental, social, or governance matters. This task creates substantial challenges. Large investors must process tens of thousands of proposals each year, and many of these proposals carry ambiguous value implications, making it difficult for shareholders to consistently choose value-maximizing actions. To manage this complexity, investors often rely on proxy advisors (PAs) that provide voting recommendations (Malenko and Shen, 2016; Shu, 2024). However, the existing literature raises persistent concerns about the quality and objectivity of these recommendations. Among them, studies document that PAs often adopt one-size-fits-all approaches (Spatt, 2021; Levit and Tsoy, 2022), face conflicts of interest (Li, 2018; Hayne and Vance, 2019; Malenko, Malenko, and Spatt, 2025), make value-destroying recommendations (Larcker, McCall, and Ormazabal, 2013, 2015; Malenko and Malenko, 2019), and biased recommendations (Ma and Xiong, 2021). Consequently, some value enhancing proposals are rejected while value destroying proposals are approved (Cuñat, Gine, and Guadalupe, 2012; Gantchev and Giannetti, 2021; Couvert, 2025). More importantly, reliance on such recommendations diminishes funds’ monitoring effectiveness and may impede the full discharge of their fiduciary duties.

Against this backdrop, we examine whether artificial intelligence (AI) can help shareholders make more value-enhancing voting decisions. Our central question is straightforward: Can AI-driven models, trained on a rich set of financial, governance, and environmental and social (E&S) data, outperform traditional proxy advisor recommendations when guiding shareholder votes? By answering this question, we address the ongoing debate over proxy advisors’ effectiveness and highlight the potential of technological innovation to improve corporate governance outcomes. Our results indicate that AI significantly outperforms the recommendations of the two leading proxy advisors, Institutional Shareholder Services (ISS) and Glass Lewis, by a substantial margin. Because shareholders and proxy advisors may not exclusively seek to maximize firm value when voting, but may also pursue other goals such as environmental or social objectives, we interpret the AI recommendations not as a prescription for how shareholders should vote, but as a value-maximizing benchmark. This benchmark allows us to quantify the economic cost

of deviating from value-maximizing voting and to study systematically when shareholders trade off firm value against other considerations.

We build a predictive framework that combines a quasi-experimental design with machine learning to generate counterfactual voting outcomes. We train a machine learning model to estimate the value implications of voting decisions, using the market reaction to the vote as a proxy for the firm’s value impact. Then, for each proposal in the out-of-sample testing set, we ask the model to predict the market reaction under two scenarios: the proposal passes, or the proposal fails. We define the AI voting recommendation as the action that yields the higher predicted reaction. Finally, we evaluate whether AI improves voting decisions by comparing market reactions for proposals where shareholders voted in line with the AI recommendation to those where they did not.

Our empirical strategy faces two significant challenges. First, proposals with a high or low ex-ante probability of passing generate little market reaction because investors have already priced in the expected outcome. Second, most votes occur during shareholder meetings, which often include several proposals as well as other announcements that can confound the market response. To address these issues, we follow the existing literature (Cuñat et al., 2012; Flammer, 2015; Ertimur, Ferri, and Oesch, 2015) and restrict the sample to proposals whose voting outcomes fall within a narrow margin around the majority threshold, in the spirit of a regression discontinuity design (RDD). These close-call proposals are contentious, and their outcomes are less likely to be anticipated by the market. Because proposals just above and just below the threshold differ only randomly, this approach allows us to identify the causal effect of proposal approval on firm value.

As briefly mentioned earlier, our first main result shows that AI-guided voting decisions systematically outperform actual shareholder behavior. Across multiple bandwidths and model specifications, proposals in which shareholders vote in line with the AI recommendation generate average cumulative abnormal returns (CARs) that are approximately 1 percentage point higher than those in which shareholders deviate from it. This pattern indicates that machine learning can identify value-enhancing voting strategies that investors frequently overlook.

Second, we compare our machine learning model’s recommendations to those of ISS, the largest proxy advisor. We regress market reactions on indicators for alignment with the ISS and the AI recommendation. Across specifications, the difference of the coeffi-

cient on the AI indicator and the one on the ISS indicator is positive and statistically significant, which shows that following the AI produces voting recommendations that are more value-enhancing-driven than ISS's. We repeat the analysis for Glass Lewis (GL), the second-largest proxy advisor, and obtain similar results. Voting outcomes that follow GL recommendations do not generate significantly higher CARs, whereas outcomes that follow the AI recommendation consistently do. These findings suggest that PAs do not make value-creating voting recommendations that our AI identifies. This result challenges the view that PAs provide the benchmark for informed voting and highlights the potential for AI-driven tools to materially improve voting outcomes.

Third, we examine which proposals the AI model supports relative to shareholders and proxy advisors. We find that, compared to AI recommendations, shareholders tend to support management proposals too frequently. In contrast, ISS and Glass Lewis are considerably more favorable toward shareholder proposals than the AI model. In particular, the AI is more likely to oppose environmental and social proposals. The model also shows less opposition to proposals at firms with weak governance or poor accounting performance, as measured by profit margins and ROA. Whereas traditional PAs seem to focus on the proposal topic to reach their recommendation, the AI model does not systematically favor or oppose specific topics, but instead incorporates a broad set of proposal and firm characteristics into its recommendations. This pattern suggests that AI recommendations are less “one-size-fits-all” and more context-dependent than those of traditional proxy advisors. Consistent with this interpretation, we show that the AI's support decisions load primarily on firm-level fundamentals, even after controlling for proposal categories, whereas proxy advisors' recommendations are largely driven by proposal type. This difference is consistent with the AI model's ability to capture non-linear interactions between proposal characteristics and firm fundamentals.

Fourth, we examine when and why shareholders vote in line with the AI recommendation. Shareholders follow the AI-recommended vote in 53% of cases, compared with 28% for ISS and 50% for GL, the two largest PAs. In other words, shareholders seem to be better able to identify value-enhancing proposals than traditional PAs. We then analyze the determinants of shareholders' alignment with the AI recommendation and uncover substantial heterogeneity. In particular, voting outcomes are more likely to align with the AI recommendation for shareholder-sponsored proposals. Importantly, we do not find

that the pursuit of environmental or social goals explains shareholders’ deviations from value-enhancing strategies.

We investigate the long-term implications of following AI recommendations. Specifically, we use our RDD approach to study the impact of voting in alignment with the AI on firm performance measures after one year. We find that when shareholders vote in alignment with the AI recommendations, target firms have higher growth in ROA, profit margin, sales, and labor productivity one year after the vote. This result indicates the positive market reaction we observe for AI-aligned votes reflects the expected firm performance improvements.

An important question is how our findings relate to the broader set of corporate voting decisions that fall outside the close-call setting. Non-contentious proposals contrast with contentious ones in that shareholders tend to agree on how to vote, as reflected in the high degree of consensus in voting outcomes. This distinction matters because it determines when AI-based recommendations are most useful and when shareholders already know how to vote. For the practical relevance of our approach, it is therefore essential to assess whether we can identify contentious proposals *ex ante*. Using the same information set available to shareholders, we show that the model predicts whether a proposal will be contentious with high accuracy. Moreover, among proposals predicted to be non-contentious, predicted outcomes align almost perfectly with realized outcomes. Taken together, these results suggest that our approach can be applied selectively, focusing on proposals for which voting decisions are less clear *ex ante*.

Our study contributes to several strands of the literature. First, we advance research on proxy voting and proxy advisors. We show that following proxy advisor recommendations does not lead to superior value outcomes. In particular, our finding that proxy advisors frequently make suboptimal recommendations provides rigorous evidence on the limitations of the existing advisory model and adds to its critiques ([Larcker et al., 2013](#), [2015](#); [Li, 2018](#); [Hayne and Vance, 2019](#); [Ma and Xiong, 2021](#); [Spatt, 2021](#); [Levit and Tsoy, 2022](#); [Malenko et al., 2025](#)). Our results also highlight the potential for technological innovation, and artificial intelligence in particular, to improve voting outcomes and enhance firm value.

Our paper is also closely related to contemporaneous work by [Lee and Souther \(2026\)](#), who use a large language model to generate voting recommendations based on proxy ad-

visor guidelines in order to identify settings in which traditional proxy advisors may face conflicts of interest. While their study focuses on proxy advisors’ objectivity relative to their stated policies, we develop a machine-learning model with the explicit objective of maximizing firm value and use it to assess the extent to which AI can help shareholders make value-enhancing voting decisions. This framework allows us to study systematically when and why shareholder voting outcomes deviate from value maximization and to quantify the associated economic consequences. In this sense, the two papers are complementary: [Lee and Souther \(2026\)](#) emphasize institutional frictions within the proxy advisory industry, whereas we focus on the implications of shareholder voting behavior for firm value.

Second, we contribute to the growing literature on the use of artificial intelligence in finance and corporate governance ([Erel, Stern, Tan, and Weisbach, 2021](#); [Van Binsbergen, Han, and Lopez-Lira, 2023](#); [Cao, Jiang, Wang, and Yang, 2024b](#); [Zhang, Zhu, and Linnainmaa, 2025](#)). [Erel et al. \(2021\)](#) use a machine learning model to identify directors who would receive higher shareholder support. While such a director-selection tool is particularly useful for nomination committees, our approach provides direct voting recommendations to shareholders across a wide range of proposals, from say-on-pay to ESG items. By using market reactions as the objective function, we do not assume that shareholders know the ex-ante value-maximizing vote; indeed, we show that they often vote suboptimally. Our regression-discontinuity design allows us to provide causal evidence on the value implications of the AI recommendations.

Finally, our paper contributes to the literature on causal machine learning by introducing a framework that combines the predictive strengths of machine learning with causal identification from a regression discontinuity design. This approach complements recent work at the intersection of machine learning and causal inference in economics and finance ([Chernozhukov, Chetverikov, Demirer, Duflo, Hansen, Newey, and Robins, 2018](#); [Athey and Imbens, 2019](#); [Wager and Athey, 2018](#); [Hansen and Siggaard, 2024](#); [Gulen, Jens, and Page, 2025](#)).

The remainder of the paper is organized as follows. Section 2 outlines the empirical methodology, the sample construction, and summary statistics. Section 3 presents the main results. Section 4 reports robustness tests and Section 5 concludes.

2 Methodology, sample description, and summary statistics

2.1 AI Voting Recommendations and Empirical Design

We develop a framework to generate AI voting recommendations for all corporate proposals. In this framework, the AI can issue only two types of recommendations: vote “For” or vote “Against.”¹ To produce these recommendations, we must specify an objective function. Several choices are possible. One option would be to recommend the vote that would maximize the voting outcome, but substantial evidence shows that shareholders often vote against value-enhancing proposals or in favor of value-reducing ones when they are uninformed (Cuñat et al., 2012; Gantchev and Giannetti, 2021; Couvert, 2025). We therefore define the objective function as maximizing firm value. Let $V_{i,1}$ and $V_{i,0}$ denote the firm value associated with proposal i passing or failing, respectively. The optimal vote is the action that maximizes firm value:

$$\text{Optimal vote}_i = \begin{cases} \textit{For} & \text{if } V_{i,1} > V_{i,0} \\ \textit{Against} & \text{otherwise} \end{cases}$$

The literature commonly uses short-window cumulative abnormal returns around voting dates to capture the market’s assessment of the value implications of voted proposals (Cuñat et al., 2012; Flammer, 2015; Ertimur et al., 2015; Gao and Huang, 2025). Following this approach, we measure the value implications of each proposal using three-day cumulative abnormal returns centered on the voting date, (CAR_{11}), computed using the Fama-French-Carhart four-factor model (Carhart, 1997).

Our identification strategy relies on close-call proposals around the majority threshold, where voting outcomes are plausibly exogenous. In this setting, proposals that narrowly pass or fail differ primarily due to quasi-random variation in voting outcomes, allowing us to identify the causal effect of proposal approval on firm value. We describe the construction of this close-call sample and provide supporting evidence for the validity of this approach in the data section below.

To generate out-of-sample predictions of the value implications of proposals under the two possible voting outcomes—passage and failure—we implement a two-stage *T-learner* architecture. In the first stage, we split the sample of proposals into a training and a testing

¹We do not allow for “Abstain,” which mirrors the recommendation structure used by proxy advisors.

set. For the training sample, we estimate two machine learning models: one using only proposals that passed ($Pass = 1$), and another using only proposals that failed ($Pass = 0$). This approach allows us to capture heterogeneous value implications depending on the voting outcome and yields two conditional expectation functions that map proposal characteristics into predicted market reactions under each outcome. Formally, we estimate:

$$\mathbb{E}[CAR_{11,i} \mid Pass_i = 1, X_i] = f_1(X_i),$$

$$\mathbb{E}[CAR_{11,i} \mid Pass_i = 0, X_i] = f_0(X_i),$$

where X_i denotes a set of observable proposal- and firm-level characteristics.

In the second stage, we use these estimates to generate counterfactual predictions for each proposal in the testing sample under both possible outcomes:

$$\widehat{CAR_{11,i}(Pass = 1)} = f_1(X_i),$$

$$\widehat{CAR_{11,i}(Pass = 0)} = f_0(X_i).$$

The underlying idea is that a proposal yields different value implications depending on whether it passes or fails the voting stage. Although only one outcome is observed in reality, these counterfactual predictions allow us to infer what the alternative outcome *could have been* according to the model. Comparing predicted values across the two counterfactual scenarios yields the model-implied value-maximizing decision. Formally, for each proposal i , the AI voting recommendation is defined as:

$$AI \text{ voting } \widehat{recommendation}_i = \begin{cases} For, & \text{if } \widehat{CAR_{11,i}(Pass = 1)} > \widehat{CAR_{11,i}(Pass = 0)}, \\ Against, & \text{otherwise.} \end{cases}$$

In essence, the model prescribes a counterfactual optimal vote, i.e., the decision that maximizes predicted CAR_{11} . This framework provides a normative benchmark for assessing observed voting behavior and highlights proposals that, according to the model, should have been voted differently. In doing so, it transforms an inherently unobservable counterfactual into a tractable and informative diagnostic of decision quality.

To preserve the temporal structure of the data and eliminate look-ahead bias, we imple-

ment a walk-forward expanding-window approach. We first use proposals from 2007–2014 as an initial training and hyperparameter-tuning period. We then generate out-of-sample predictions sequentially for 2014–2024, re-estimating the model each year using only information available up to that point. Following this procedure, we use the period 2015–2019 to construct the stacked ensemble and reserve the final period, 2020–2024, for model evaluation, ensuring that this period is not used in model estimation or tuning. Because all predictions are based solely on past data, this procedure mimics a real-time decision environment.

Finally, we evaluate whether the model-implied voting strategy enhances shareholder value by comparing proposals where shareholders’ actual votes align with the AI recommendation to those where they do not. Specifically, we estimate:

$$\begin{aligned}\Delta\text{CAR}_{11} = & \mathbb{E}[\text{CAR}_{11} \mid \text{actual vote} = \text{AI voting recommendation}] \\ & - \mathbb{E}[\text{CAR}_{11} \mid \text{actual vote} \neq \text{AI voting recommendation}]\end{aligned}$$

A positive and statistically significant ΔCAR_{11} indicates that the model-guided voting strategy outperforms actual shareholder behavior.

2.2 Machine Learning Setup

Previous studies suggest that machine learning algorithms can capture complex, non-linear relationships among variables, which may be useful for predicting financial returns in settings where traditional models struggle to represent such patterns (Goodfellow, Bengio, and Courville, 2016; James, Witten, Hastie, Tibshirani, and Taylor, 2023). Following Cao, Jiang, Wang, and Yang (2024a), we construct an ensemble with three non-linear base learners: Extreme Gradient Boosting (XGB), Random Forest (RF), and a Neural Network (NN). These learning models offer complementary strengths. Extreme Gradient Boosting is effective for structured tabular data and incorporates regularization to mitigate overfitting. Random Forest delivers robust predictions and interpretability through ensemble averaging. Neural Networks are particularly well-suited for capturing highly non-linear and hierarchical patterns in the data.

In the first stage of the estimation procedure, we train each of the three base learners separately within each treatment state, that is, on proposals that pass and proposals that

fail. This yields a set of predicted counterfactual outcomes for each proposal.

In the second stage, we construct our stacked ensemble by combining the predictions of the base learners. Specifically, we train an XGB meta-learner on the base learners' out-of-sample predictions to produce final estimates of CAR_{11} under each voting outcome. This stacking procedure allows us to aggregate information across models and improve out-of-sample prediction performance.

2.3 Sample Construction

2.3.1 Data sources

We obtain data on proposals from the Company Voting Results database within the ISS Voting Analytics suite. The database covers all shareholder and management proposals at US firms from 2007 to 2024, yielding a sample of more than 200,000 proposals. It provides comprehensive information on proposal topics, sponsor types, and final voting outcomes. In addition to proposal data, we obtain firms' financial and accounting information from the CRSP/Compustat Merged dataset, and stock prices from CRSP. We source institutional ownership data from Thomson Reuters (13F filings), incorporate environmental, social, and governance (ESG) metrics from MSCI ESG Ratings, and add analysts' recommendations from IBES.

An often-voiced concern of machine-learning applications in finance is the look-ahead bias (Zhang et al., 2025), which arises when researchers use information that was not available at the time predictions occur. To avoid this bias, we adopt a conservative merging strategy and ensure that all firm-level variables are observable at the time of the vote. In practice, this implies that for each proposal we use only information from the fiscal period preceding the voting date.

Finally, because proxy advisor recommendations are not directly available, we incorporate recommendations for both ISS and Glass Lewis using the imputed data constructed by Zytznick (2025), which infers recommendations from publicly available institutional voting records.

2.3.2 Constructing the AI information set

We provide the machine learning model with a set of proposal- and firm-level variables that are observable to shareholders and proxy advisors at the time of the vote. Our objective

is not to supply the model with additional information, but rather to assess whether it can extract value implications more effectively from commonly available data.

To limit the risk of overfitting, we rely on a parsimonious set of variables capturing key dimensions of proposals and firm characteristics. On the proposal side, we include an indicator for whether the proposal is shareholder-sponsored, as voting behavior and value implications often differ between shareholder and management proposals. We also account for proposal content by grouping proposals into a small number of topic categories. These categories capture broad dimensions such as governance, compensation, capital structure, and ESG issues (details on the classification procedure are provided in the appendix).

On the firm side, we include variables designed to capture financial performance, governance structure, valuation, and the firm’s information environment. Profitability measures such as ROA and profit margin proxy for operating performance, while one-year stock returns reflect recent market performance. Market capitalization, book leverage, and the book-to-market ratio capture firm size, capital structure, and valuation. We include dividend yield to reflect payout policy. ESG pillar scores from MSCI provide information on governance quality and sustainability performance. Finally, analyst recommendations from IBES proxy for market expectations and the firm’s information environment.

We restrict our sample to proposals whose voting outcomes fall within a 20% margin around the majority threshold. This restriction serves two purposes. First, proposals with very high or very low ex-ante approval probabilities generate limited market reactions because their outcomes are largely anticipated, and we follow [Gao and Huang \(2025\)](#) in adopting this convention. Second, focusing on close-call proposals allows us to exploit quasi-random variation in proposal passage around the voting threshold. This approach follows a well-established literature using regression discontinuity designs in shareholder voting settings ([Cuñat et al., 2012](#); [Flammer, 2015](#)) and enables us to identify the causal effect of proposal approval on firm value, measured by short-window CARs.

For this design to yield credible causal inference, proposals must pass in a locally random manner near the threshold. We assess this assumption using a McCrary density test ([McCrary, 2008](#)). Figure 1 reports the estimated density around the cutoff. The absence of a discontinuity, together with a p-value of 0.79, indicates no evidence of manipulation and supports the validity of the design.

2.4 Summary statistics

Table 1 presents descriptive statistics for the proposal- and firm-level variables in our close-call sample, defined as proposals receiving 30%–70% shareholder support. Panel A documents substantial heterogeneity across proposal types and voting inputs. The mean Pass indicator equals 0.59, while the average votes-for ratio is 0.52, mechanically reflecting the sample’s focus on contentious items. Shareholder-sponsored proposals represent 41% of observations. Management issues a “For” recommendation in 59% of cases, whereas ISS and Glass Lewis recommend support in 48% and 57% of cases, respectively. Topic classifications show that close-call items concentrate in Board & Governance (41%) and Compensation (24%), with ESG and Articles/Bylaws proposals each representing roughly 11% of the sample.

Panel B reports firm-year-level characteristics for the underlying firms. The median firm-year observation has total assets of \$8.75 billion and sales of \$4.31 billion. Institutional ownership averages 81%, and the median book leverage is 0.31. Profitability measures show a median ROA of 4% and a median profit margin of 8%. The median Tobin’s Q equals 1.72, and the median book-to-market ratio equals 1.04. ESG pillar scores vary across firms, with mean values of 5.12 (governance), 5.17 (environmental), and 4.31 (social).

3 Results

3.1 AI vs. Shareholders?

First, we examine whether our AI recommendations support voting decisions that maximize firm value. To do so, we compare market reactions to proposal votes in cases where the voting outcome aligns with our AI recommendation and in cases where it does not. Table 2 reports our first main results across multiple bandwidths. Columns 1 to 4 restrict the sample to close-call proposals whose voting outcomes fall within 10%, 7.5%, 5%, and 2.5% of the majority threshold. Column 5 restricts the sample to proposals that fall within the optimal bandwidth around the threshold, which we compute following [Calonico, Cattaneo, and Farrell \(2020\)](#) and which equals 3.8%. Column 6 presents the results for the full model, which includes all proposals that fall within the 20% margin around the majority threshold. For Column 6, we include control variables as well as year and proposal-topic fixed effects. For all estimates, we cluster standard errors at the firm level.

AI – aligned votes denotes proposals for which the voting outcome aligns with the AI recommendation. *AI – non – aligned votes* denotes proposals for which the voting outcome differs from the AI recommendation. *Difference* reports the gap in average CARs between the two subsamples and provides a one-sided test of whether *AI – aligned votes* produce larger market reactions than *AI – non – aligned votes*. Across multiple bandwidths around the voting threshold, the results consistently show that votes aligned with AI recommendations generate higher CARs than votes that do not. The absolute difference is positive and statistically significant for every bandwidth. On average, the difference in CARs across specifications equals 0.97%. The gap widens as the bandwidth narrows, with the strongest market reaction at the narrowest threshold, which reaches 1.51%. Overall, the evidence indicates that votes whose outcome is aligned with our AI model generate the largest market reactions, thereby suggesting that the AI model, on average, identifies the firm-value-maximizing voting strategies.

To further illustrate the ability of AI to identify the optimal voting policy, we focus on proposals whose predicted difference in market reactions between the two counterfactual outcomes exceeds the median value, as not all proposals have value implications. For proposals within a 2.5% bandwidth of the majority threshold, we then form two portfolios: one containing proposals whose voting outcomes align with the AI recommendation and one containing proposals whose outcomes contrast with the AI recommendation. Figure 2 plots long-run CARs for these two portfolios. The trajectories show that market reactions are substantially more favorable when shareholders follow the AI’s voting strategy, and the divergence persists for up to 180 trading days after the vote. In terms of magnitude, the spread between the two portfolios reaches 13 percentage points over the 180-day window. This pattern suggests that AI-guided voting not only enhances short-term value creation but also produces lasting effects on firm performance.

Overall, our results show that AI models trained on comprehensive financial, governance, and market data can systematically identify value-enhancing voting decisions that shareholders often overlook. The statistically significant and economically meaningful CAR differentials between AI-aligned and AI-non-aligned votes provide strong evidence that machine learning can improve corporate governance outcomes. These findings carry important implications for institutional investors, as they suggest that integrating AI-driven decision-support tools into the proxy voting process can mitigate suboptimal

voting and enhance shareholder value.

3.2 AI vs. Proxy advisors

Our results so far indicate that AI can generate voting recommendations that outperform the aggregate decisions of shareholders. However, many shareholders, particularly institutional investors with fiduciary duties to their clients, rely on proxy voting advisors. Because these advisors present themselves as experts in formulating voting recommendations, one might expect their guidance to be associated with positive value implications. At the same time, proxy advisors have faced substantial criticism. Common concerns include that they issue one-size-fits-all recommendations that may harm firm value, that they face conflicts of interest when they sell consulting services to the same firms for which they provide voting advice, and that their recommendations may reflect ideological positions. As a result, it is not obvious that their recommendations lead to higher value creation.

To examine this question, we apply our analytical framework to proxy advisor recommendations. Specifically, we compare stock market reactions to votes on proposals whose outcomes align with a proxy advisor’s recommendation with those whose outcomes deviate from it. We report the results in Table 3.

Panel A of Table 3 focuses on the largest proxy advisor, ISS. *ISS – aligned votes* is the sample of proposals whose voting result aligns with the ISS voting recommendation. *ISS – non – aligned votes* is the sample of proposals whose voting result does not align with the ISS recommendation. *Difference* reports the average difference in market reactions between the two samples across the different thresholds. It appears that this difference is negative and not statistically significant for almost all thresholds. This result suggests that the recommendations issued by ISS do not lead to higher firm value implications

Moreover, the second part of the table compares the ability of AI to identify value-enhancing voting strategies relative to ISS. To do so, we estimate the following model:

$$\Delta\text{CAR}_{11} = \text{AI-aligned} + \text{ISS-aligned} + \Gamma'\text{Controls} + \delta_t + \delta_c + \epsilon$$

where *AI – aligned* is an indicator variable that takes the value 1 if the voting outcome

aligns with the AI recommendation and 0 otherwise. *ISS-aligned* is an indicator variable that takes the value 1 if the voting outcome aligns with the ISS recommendation, and 0 otherwise. Column 6 includes control variables as well as year and proposal-topic fixed effects, δ_t and δ_c , respectively.

It appears that for all margins except for the widest one of 20%, *AI-aligned* is statistically significant and positively associated with the three-day CAR. In contrast, *ISS-aligned* is not statistically significant for any of the thresholds. We then examine whether the coefficients on *AI-aligned* are larger than on *ISS-aligned*. *Difference* reports the results of the test. For all margins, the difference is positive, and it is statistically significant for all except 20% margin. On average, alignment with AI recommendations is associated with a 1% higher CARs than alignment with ISS. This result implies that the recommendations generated by our AI model outperform those issued by ISS, the largest proxy advisor.

We then turn to the second largest proxy advisor, Glass Lewis, and present the results in Panel B of Table 3. Similar to what we find for ISS, we do not find that proposals whose voting result is aligned with GL’s recommendation generate a market reaction that is significantly different from the market reaction of proposals whose voting result is not aligned with GL’s recommendation. This finding suggests that, similar to ISS, the recommendations made by GL are not associated with higher value implications. Furthermore, we compare the AI recommendations to the GL ones. We find that being aligned with AI recommendations is positively associated with market reactions to proposal votes. In contrast, being aligned with Glass Lewis recommendations does not have any statistically significant association with the market reaction. Finally, we analyze whether the *AI-aligned* coefficients are larger than the *GL-aligned* coefficients. We find that the difference is positive and statistically significant for all the margins. On average, alignment with AI recommendations is associated with a 0.89% higher CARs than alignment with ISS.

Overall, our results show that the AI model’s proposed voting strategy outperforms both ISS and GL recommendations in terms of firm value implications following the voting stage. The economic magnitude of the CAR differentials is substantial, which suggests that AI models can systematically identify value-enhancing voting decisions that proxy advisors frequently miss. This evidence challenges the conventional view that proxy ad-

visors represent the gold standard in shareholder voting. Instead, our findings indicate that AI-driven decision-support tools can materially improve voting outcomes and generate tangible benefits for firm value. For institutional investors, integrating AI into the proxy voting process may therefore offer a promising way to mitigate the limitations of traditional advisory services and strengthen the effectiveness of corporate governance.

3.3 AI+proxy advisors vs. AI

So far, we have shown that a machine learning model can generate voting recommendations that outperform both shareholders and the two main proxy advisors in identifying voting strategies that maximize firm value implications. The information set that we provide to our model consists of publicly available data. Although our model performs better than proxy advisors in producing value-enhancing voting recommendations, this does not imply that proxy advisors fail to generate information that could improve the model’s predictions. Proxy advisors may have access to public information that we do not include in our information set. They may also have access to non-public or soft information, for example, through their interactions with firms. Finally, they may process these pieces of information in a different way.

To investigate the value of proxy advisor recommendations for our AI model, we sequentially add the recommendations of the two proxy advisors to the model’s information set and retrain the model. We begin with the voting recommendations of the largest proxy advisor, ISS. After training the model, we apply our empirical strategy to obtain AI recommendations and compare the model’s ability to generate value-enhancing guidance with and without the ISS recommendations included in its information set. We present the results in Panel A of Table 4.

$AI - aligned$ is an indicator variable that takes the value 1 when shareholders vote in alignment with our benchmark AI recommendations, and 0 otherwise. $AI + ISS - aligned$ is an indicator variable that takes the value one when shareholders vote in alignment with the AI recommendations when the AI information set includes the ISS recommendations. *Difference* is the difference between the two coefficients.

We find that the difference coefficient is positive but not statistically significant at the 10% level for almost all the margins. In other words, we do not find any evidence that the AI model performs better when it is not provided with the ISS recommendations.

e then turn to GL and perform a similar test. We present the results in Panel B of Table 4. As in the ISS case, we do not find evidence that the AI model performs better for any of the thresholds when the Glass Lewis recommendations are included in its information set. In fact, we find that the difference between the *AI – aligned* and the *AI + GL – aligned* coefficients is always positive and statistically significant for all bandwidths except the widest one. Although this finding may appear surprising, it is consistent with the possibility that the GL recommendations introduce noise that leads the model to overfit.

Overall, our results are consistent with proxy advisors not providing additional information that improves the performance of our AI model. This finding aligns with common critiques of proxy advisors, particularly the concern that they rely on a one-size-fits-all approach when formulating their recommendations.

3.4 Which proposals does the AI support?

To shed light on the mechanisms underlying the model’s proposed voting strategy, we begin by examining the distribution of support across shareholders, the AI model, and traditional proxy advisors. Table 5 reports support rates for all proposals and distinguishes between shareholder-sponsored and management-sponsored proposals. The figures reveal substantial heterogeneity in voting behavior across decision agents.

From Column 1, it appears that shareholders support 56% of all proposals. However, support rates differ sharply by sponsorship. Shareholders support shareholder-sponsored proposals in only 16% of cases, while they support management-sponsored proposals in 87% of cases. In contrast, the model’s voting strategy in Column 2 supports only 22% of all proposals, regardless of sponsorship. The results show that the model supports shareholder-sponsored proposals is low, 12% of cases. Moreover, its support to management-sponsored proposals is much lower than shareholders’ support to these proposals, only 23%. Column 3 presents the support rates for ISS. Across all proposals and irrespective of sponsorship, ISS supports 51% of proposals. Taking sponsorship into account, ISS supports 91% of shareholder-sponsored proposals and only 20% of management-sponsored proposals, which reflects a markedly more activist stance than the model’s support rates. Glass Lewis adopts a somewhat more moderate position, recommending support for 58% of all proposals, including 69% of shareholder proposals and

49% of management proposals.

Taken together, these results show that shareholders remain largely aligned with management, while proxy advisors, especially ISS, tend to favor shareholder-sponsored initiatives strongly. The AI-model-based results occupy a more middle-ground position, neither supporting an excessive share of shareholder-sponsored proposals nor an excessive share of management-sponsored proposals. This pattern reflects a more differentiated support profile that does not systematically privilege either side. These findings suggest that the model may internalize elements of both shareholder and proxy advisor preferences while maintaining a more balanced assessment of proposal merits, rather than relying primarily on proposal sponsorship.

To further shed light on the determinants of proposal support across shareholders, the AI model, ISS, and GL, we estimate a set of regressions designed to identify which factors influence support decisions for each of the four agents. Table 6 reports the results. Models (1)–(4) correspond to separate regressions estimated for each voting strategy, where the dependent variable equals 1 if the respective decision-maker supports a proposal, and 0 otherwise. All specifications include year and proposal-topic fixed effects, and standard errors are clustered at the firm level.

The results reveal several notable patterns regarding how different decision-makers form their support decisions. First, proposal characteristics play a central role for shareholders and, especially, for the two traditional proxy advisors. Shareholders are substantially less likely to support shareholder-sponsored proposals, whereas ISS and Glass Lewis are significantly more supportive of these proposals relative to management-sponsored items. Moreover, the magnitude of the coefficients on proposal categories is large for ISS and GL, consistent with these advisors relying heavily on proposal type as a primary input into their recommendation policies. In contrast, the AI model places only limited weight on proposal characteristics. With the exception of a negative coefficient on ESG proposals, indicating that the AI is somewhat more skeptical toward environmental and social items, the proposal topic plays only a minor role in shaping the model’s support decisions.

Second, the contrast in how firm characteristics enter the support rules is striking. The AI model assigns economically meaningful weight to several firm-level variables: for example, it reduces support for proposals at firms with higher governance scores, and it exhibits differential support based on profitability and valuation metrics. Shareholders also

respond to selected firm characteristics, though to a lesser extent. In sharp contrast, ISS incorporates only a subset of firm variables into its recommendations, and GL does not load significantly on any firm characteristic. These results are consistent with the often-voiced concern that traditional proxy advisors adopt a one-size-fits-all voting recommendation. The AI model conditions its recommendations on a broader and more granular assessment of the firm’s underlying fundamentals. These results are consistent with the often-voiced concern that traditional proxy advisors adopt a one-size-fits-all voting recommendation.

Third, the differences in model fit provide additional insight. The AI specification exhibits a substantially lower adjusted R^2 relative to the shareholder and ISS regressions. This pattern is consistent with the fact that the AI recommendations arise from a non-linear learning process that captures complex interactions across proposal and firm attributes, relationships that are not well approximated by a linear probability model. Glass Lewis also displays a comparatively low adjusted R^2 , which may reflect a similarly non-linear or rule-based internal approach, though one that, unlike the AI, relies almost exclusively on proposal-level criteria rather than firm fundamentals.

3.5 When do shareholders align with AI recommendations?

A central question in evaluating the potential of artificial intelligence in proxy voting is identifying the circumstances under which shareholders’ actual votes align with the AI-recommended outcome. This section examines both the unconditional rates of AI-aligned voting and the determinants that drive such alignment.

Table ?? reports the percentage of proposals for which shareholder voting outcomes align with the recommendations of the AI model (Column 1), ISS (Column 2), and Glass Lewis (Column 3). Shareholders follow the AI recommendation in 53.4% of proposals, compared with 28.1% alignment with ISS and 50.3% alignment with Glass Lewis. Thus, shareholders are more aligned with the AI than with either proxy advisor, suggesting that shareholders may, in fact, do a better job than ISS or Glass Lewis at identifying value-enhancing voting decisions. Alignment with the AI differs substantially across proposal types: shareholders follow the model in 77.8% of shareholder-sponsored proposals but only 34.7% of management-sponsored proposals. This pattern indicates that shareholders are considerably more likely to vote in accordance with the AI’s value-maximizing recommendations when evaluating shareholder-initiated proposals.

To further examine the mechanisms underlying voting in alignment with AI, Table ?? presents regression results in which the dependent variable equals one if the shareholder’s vote aligns with the AI model’s recommendation. Columns 2 and 3 report comparable specifications for alignment with ISS and Glass Lewis recommendations. All regressions include proposal-topic and year fixed effects, and standard errors are clustered at the firm level.

Shareholder-sponsored proposals are substantially more likely to be voted in accordance with the AI recommendation, mirroring the unconditional alignment patterns. Proposal characteristics more generally are strong predictors of alignment with ISS and GL, where several topic indicators load with economically large and statistically significant coefficients. By contrast, alignment with the AI shows little sensitivity to proposal categories, and, importantly, the ESG topic indicator is not statistically significant. Thus, pursuing ESG goals does not appear to explain deviations from the AI’s value-maximizing recommendations.

Firm-level characteristics further differentiate the three alignment patterns. Alignment with the AI varies systematically with several firm attributes: for example, lower alignment at firms with higher social pillar scores and higher alignment at firms with greater dividend yields. Alignment with ISS and GL also reflects some firm characteristics. But for GL, the set of significant firm-level coefficients is narrower than in the AI specification. Overall, these results indicate that, relative to proxy-advisor-based alignments, AI-consistent alignment is more firmly grounded in firm fundamentals rather than proposal labels.

3.6 How would firm structure change if shareholders followed AI recommendations?

A natural question is how firm-level governance structures would change if shareholders fully followed the AI’s voting recommendations. The preceding sections establish that the AI model identifies value-enhancing voting strategies and that shareholders do not always follow these strategies. The counterfactual exercise in this section evaluates the governance consequences of full adherence to the AI’s recommendations. Because the proxy voting process directly determines key governance outcomes, systematic differences between actual and AI-consistent votes can translate into meaningful structural changes.

Table 9 reports this counterfactual analysis by comparing actual shareholder voting outcomes with the outcomes that would arise if shareholders voted in full alignment with the AI. The table displays, for the ten most common proposal topics, the distribution of proposals according to (i) whether the AI recommends voting For or Against, and (ii) whether shareholders currently follow or deviate from the AI recommendation. All values are expressed as percentages, and topics are sorted by frequency in the sample. Because the analysis is restricted to proposals within a ± 20 -percentage-point margin of the majority threshold, management proposals in this sample are unusually contested relative to their near-unanimous pass rates in the broader population (e.g., director elections receive on average 95% approval), whereas shareholder proposals commonly attract far lower support even outside the close-call window.

The counterfactual distributions in Table 9 suggest that full compliance with the AI's recommendations would lead to economically meaningful shifts in firm structure. The largest effects arise for management-sponsored governance proposals, including director elections and say-on-pay, where the AI frequently recommends voting Against. Given that these items in the restricted sample are already unusually contested for management proposals, full adherence to the AI would imply a substantial reallocation of votes away from management, leading to more frequent challenges to incumbent directors and to executive compensation practices. In the broader setting where such proposals usually pass overwhelmingly, these counterfactual changes translate into meaningful governance adjustments.

The counterfactual effects are also sizable for several shareholder-sponsored governance proposals, notably CEO-chairman separation, proxy access, and written consent. For CEO-chairman separation, the AI often recommends voting For, yet shareholders frequently vote Against those recommendations; full alignment with the AI would therefore materially increase support for independent board leadership. At the same time, the AI does not recommend separation uniformly, indicating that its guidance is not mechanical but conditions on firm-specific characteristics, consistent with a nonlinear assessment of when separation enhances value. Similar asymmetries appear for proxy access and written consent: when the AI recommends voting For, shareholders often vote Against; under full compliance, support would shift materially toward these mechanisms, expanding shareholder nomination rights and the ability to act outside annual meetings.

These patterns underscore a central feature of the AI’s decision rule: recommendations are not uniform within a topic. Even for proposals that shareholder advocates frequently support, such as CEO-chairman separation, proxy access, and written consent, the AI does not always recommend voting For. The distribution of proposals across the four outcome cells in Table 9 indicates that the AI conditions on firm- and proposal-specific characteristics rather than applying a simple governance checklist. Consequently, the counterfactual governance outcomes reflect not only a shift toward stronger shareholder rights but also a more selective, context-specific application of those rights, consistent with value maximization rather than blanket activism.

Overall, Table 9 indicates that full alignment with AI recommendations would lead to meaningful changes in firm structure, particularly in areas related to board independence, executive compensation, and green issues. While magnitudes vary across topics, a consistent pattern emerges: the AI identifies targeted proposals, often those already facing low support within the close-call window, where governance reforms are most likely to enhance firm value.

3.7 Long-term impacts of following AI recommendations

Next, we analyze the long-term implications of following AI recommendations. To this end, we explore whether AI-aligned votes are associated with subsequent improvements in operating performance, in particular year-over-year growth in ROA, profit margin, ROE, Tobin’s Q, sales, labor productivity, and capex, compared to AI-non-aligned votes. Using our RDD setting, we regress these outcomes onto our *AI – aligned* variable, which takes the value 1 if the voting result aligns with the AI recommendation, and 0 otherwise. For all specifications, we include firm and year fixed effects, and standard errors are clustered at the firm level. Panel A evaluates the fiscal year of the vote and serves as a placebo for immediate operational changes; Panel B evaluates the subsequent fiscal year to capture post-vote adjustments. The results are presented in Table 10.

In Panel A, the coefficients on AI-aligned votes are small and statistically indistinguishable from zero across all outcomes, consistent with limited scope for contemporaneous operational shifts around the meeting date. In Panel B, firms with AI-aligned votes exhibit significantly higher growth in ROA and profit margin, alongside significantly faster growth in sales and labor productivity; effects on ROE, Tobin’s Q, and capex are not statistically

significant. Together, these results indicate that following the AI’s value-maximizing voting recommendations is associated with genuine subsequent operating gains rather than transitory price effects.

3.8 Non-contentious proposals

Our empirical design focuses on proposals whose voting outcomes fall within a narrow margin around the majority threshold. While this restriction is necessary for causal identification, it raises an important question: can investors identify, *ex ante*, which proposals are likely to fall into this category? This question is central for two reasons. First, if contentious proposals cannot be anticipated beforehand, the relevance of our results for real-world decision-making would be limited. Second, the ability to predict contentiousness allows us to distinguish between proposals that are trivial to evaluate and those that present real challenges to shareholders. For non-contentious proposals, the optimal voting strategy is largely transparent, as reflected in the high degree of consensus among shareholders. In contrast, contentious proposals are characterized by disagreement and uncertainty about value implications. As a result, predicting contentiousness is key to understanding where AI-based decision tools are most valuable.

We evaluate the model’s ability to distinguish between proposals that are contentious and those that are not, training our machine learning model to predict voting outcomes rather than market reactions. We use the same ensemble learning framework and the same set of independent variables as in our main specification, but replace the dependent variable with voting outcomes. Because the objective is to predict voting results, we include the full population of proposals rather than restricting the analysis to those within a 20% margin of the majority threshold.

We report the results in Panel A of Table [IA1](#). The model predicts whether a proposal is contentious with an out-of-sample accuracy of 97.4%, and an area under the ROC curve of 96.5%, as shown in Figure [IA1](#). These results indicate that the model is highly effective in distinguishing between proposals that are *ex ante* straightforward and those that are likely to generate disagreement among shareholders.

We next examine voting outcomes among proposals that are predicted to be non-contentious. Panel B of Table [IA1](#) shows that, within this subset, proposals with predicted support above 70% pass 99.9% of the time, while those with predicted support below 30%

fail 97.4% of the time. The overall out-of-sample accuracy is 99.8%, with an area under the ROC curve of 99.8%. These results confirm that the model can nearly perfectly predict the outcome of proposals that it classifies as non-contentious.

Taken together, these findings point to a simple but important economic interpretation. Shareholder voting problems are highly heterogeneous. For the majority of proposals, the optimal voting decision is largely transparent, and shareholders coordinate on the outcome, leading to predictable voting patterns. In contrast, a smaller subset of proposals, those that are ex ante contentious, gives rise to genuine disagreement and uncertainty about value implications. These are precisely the settings in which voting outcomes are difficult to anticipate and where informational frictions are most likely to affect shareholder decisions.

4 Robustness tests

4.1 Robustness test 1 - Excluding cases with confounding proposals

Although in most cases there is only one proposal whose voting result falls within a 20% margin around the majority threshold, it is possible for several proposals at the same annual general meeting to have close-call outcomes. This situation poses a challenge for our methodology when some of these close-call proposals follow the AI recommendation while others do not. In such cases, the estimated effects may be downward-biased because we cannot disentangle the value implications of each proposal when only a single market reaction is observed. To address this concern, we conduct a robustness test in which we exclude all instances where multiple close-call proposals do not either all follow the AI recommendation or all deviate from it. The results of this test are presented in Table [IA2](#) of the internet appendix and show that our findings are robust to restricting the analysis to this special subsample. The difference in market reaction between cases in which shareholders voted in alignment with the AI model and cases in which they did not is positive and statistically significant for each margin. The effect remains positive and significant for all bandwidths. In terms of magnitude, AI-aligned votes are associated with an average 1.2% CARs, compared to 1% when including all proposals.

4.2 Robustness test 2 - Predicted contentious proposals

Our main empirical strategy focuses on proposals whose voting outcomes fall within a narrow margin around the majority threshold. A natural concern is that shareholders do not know *ex ante* which proposals will fall into this category. In Section 3.8, we show that the model can accurately predict whether proposals are contentious and that voting outcomes for non-contentious proposals are highly predictable.

Building on this insight, we construct an alternative close-call sample based on predicted voting outcomes rather than realized ones. Specifically, we train the model to predict voting results using the full population of proposals and retain those whose *predicted* vote shares fall within a 20% margin of the majority threshold. We then compute AI voting recommendations within this predicted close-call sample.

Table 11 presents the results. Across all thresholds, the difference between the CARs for AI-aligned and AI-non-aligned votes remains positive and statistically significant for most specifications. These findings confirm that our results are robust when the subset of contentious proposals is identified *ex ante* rather than *ex post*.

5 Conclusion

This paper provides causal evidence that artificial intelligence can materially improve shareholder voting outcomes. By combining regression discontinuity with ensemble machine learning, we show that AI-generated recommendations consistently outperform both actual shareholder behavior and the guidance of traditional proxy advisors. Strikingly, shareholders are more likely to vote in alignment with the AI model than with either ISS or Glass Lewis, which highlights that investors often recognize when not to rely on proxy advisors in order to make value-enhancing voting decisions. At the same time, proxy advisor recommendations do not appear to enhance the predictive power of the AI model, suggesting that the information content of traditional advisory guidance is largely redundant or even misleading.

Our analysis also reveals why shareholders often fail to follow the AI-optimal vote: they exhibit excessive support for management-sponsored proposals, even when such proposals reduce firm value. In contrast, the AI model incorporates a richer set of proposal- and firm-level characteristics, exploiting non-linear interactions that traditional advisors overlook.

This methodological advantage allows AI to identify value-maximizing strategies across diverse proposal types without resorting to one-size-fits-all heuristics.

Finally, we document that the benefits of AI-guided voting extend beyond short-term market reactions. Firms whose proposals are decided in line with AI recommendations experience higher profitability, productivity, and sales growth in the subsequent year. These long-term improvements confirm that the abnormal returns associated with AI-aligned votes reflect genuine enhancements in firm performance rather than temporary price effects. With regard to non-contentious proposals, we show that voting outcomes exhibit a high degree of consensus and can be predicted with little ambiguity. Because such cases can be identified *ex ante*, this allows our approach to be focused on proposals where voting decisions are less clear. Taken together, our findings highlight the potential for AI to reshape corporate governance by reducing reliance on conflicted proxy advisors, correcting shareholder biases, and fostering sustained value creation.

References

- Athey, S. and G. W. Imbens (2019). Machine learning methods that economists should know about. *Annual Review of Economics* 11(1), 685–725.
- Calonico, S., M. D. Cattaneo, and M. H. Farrell (2020). Optimal bandwidth choice for robust bias-corrected inference in regression discontinuity designs. *The Econometrics Journal* 23(2), 192–210.
- Cao, S., W. Jiang, J. Wang, and B. Yang (2024a). From man vs. machine to man + machine: The art and ai of stock analyses. *Journal of Financial Economics* 160, 1–22.
- Cao, S., W. Jiang, J. Wang, and B. Yang (2024b). From man vs. machine to man+ machine: The art and ai of stock analyses. *Journal of Financial Economics* 160, 103910.
- Carhart, M. M. (1997). On persistence in mutual fund performance. *Journal of finance* 52(1), 57–82.
- Chernozhukov, V., D. Chetverikov, M. Demirer, E. Duflo, C. Hansen, W. Newey, and J. Robins (2018). Double/debiased machine learning for treatment and structural parameters.
- Couvert, M. (2025). What are the firm value implications of sec-challenged shareholder proposals? *Management Science* 71(6), 4728–4756.
- Cuñat, V., M. Gine, and M. Guadalupe (2012). The vote is cast: The effect of corporate governance on shareholder value. *The Journal of Finance* 67(5), 1943–1977.
- Erel, I., H. Stern, Léa C. Tan, and S. Weisbach, Michael (2021, April). Selecting directors using machine learning. *The Review of Financial Studies* 34(7), 3226–3264.
- Ertimur, Y., F. Ferri, and D. Oesch (2015). Does the director election system matter? Evidence from majority voting. *Review of Accounting Studies* 20(1), 1–41.
- Flammer, C. (2015). Does corporate social responsibility lead to superior financial performance? a regression discontinuity approach. *Management science* 61(11), 2549–2568.
- Gantchev, N. and M. Giannetti (2021). The costs and benefits of shareholder democracy: Gadflies and low-cost activism. *The Review of Financial Studies* 34(12), 5629–5675.
- Gao, M. and J. Huang (2025). Informed voting. *The Review of Financial Studies* 38(4), 1167–1210.
- Goodfellow, I., Y. Bengio, and A. Courville (2016). *Deep Learning*. Cambridge, MA: MIT Press.
- Gulen, H., C. E. Jens, and T. B. Page (2025). Balancing external vs. internal validity: An application of causal forest in finance. *Management Science*.
- Hansen, J. H. and M. V. Siggaard (2024). Double machine learning: Explaining the post-earnings announcement drift. *Journal of Financial and Quantitative Analysis* 59(3), 1003–1030.
- Hayne, C. and M. Vance (2019). Information intermediary or de facto standard setter? field evidence on the indirect and direct influence of proxy advisors. *Journal of Accounting Research* 57(4), 969–1011.

- James, G., D. Witten, T. Hastie, R. Tibshirani, and J. Taylor (2023). *An Introduction to Statistical Learning: With Applications in Python* (2nd ed.). New York: Springer.
- Larcker, D. F., A. L. McCall, and G. Ormazabal (2013). Proxy advisory firms and stock option repricing. *Journal of Accounting and Economics* 56(2-3), 149–169.
- Larcker, D. F., A. L. McCall, and G. Ormazabal (2015). Outsourcing shareholder voting to proxy advisory firms. *The Journal of Law and Economics* 58(1), 173–204.
- Lee, C. and M. E. Souther (2026). Beyond bias: Ai as a proxy advisor. *European Corporate Governance Institute–Finance Working Paper Forthcoming*.
- Levit, D. and A. Tsoy (2022). A theory of one-size-fits-all recommendations. *American Economic Journal: Microeconomics* 14(4), 318–347.
- Li, T. (2018). Outsourcing corporate governance: Conflicts of interest within the proxy advisory industry. *Management Science* 64(6), 2951–2971.
- Ma, S. and Y. Xiong (2021). Information bias in the proxy advisory market. *The Review of Corporate Finance Studies* 10(1), 82–135.
- Malenko, A. and N. Malenko (2019). Proxy advisory firms: The economics of selling information to voters. *The Journal of Finance* 74(5), 2441–2490.
- Malenko, A., N. Malenko, and C. Spatt (2025). Creating controversy in proxy voting advice. *The Journal of Finance* 80(4), 2303–2354.
- Malenko, N. and Y. Shen (2016). The role of proxy advisory firms: Evidence from a regression-discontinuity design. *The Review of Financial Studies* 29(12), 3394–3427.
- McCrary, J. (2008). Manipulation of the running variable in the regression discontinuity design: A density test. *Journal of econometrics* 142(2), 698–714.
- Shu, C. (2024). The proxy advisory industry: Influencing and being influenced. *Journal of Financial Economics* 154, 103810.
- Spatt, C. S. (2021). Proxy advisory firms, governance, market failure, and regulation. *The Review of Corporate Finance Studies* 10(1), 136–157.
- Van Binsbergen, J. H., X. Han, and A. Lopez-Lira (2023). Man versus machine learning: The term structure of earnings expectations and conditional biases. *The Review of financial studies* 36(6), 2361–2396.
- Wager, S. and S. Athey (2018). Estimation and inference of heterogeneous treatment effects using random forests. *Journal of the American Statistical Association* 113(523), 1228–1242.
- Zhang, Y., Y. Zhu, and J. T. Linnainmaa (2025). Man versus machine learning revisited. *The Review of Financial Studies* 38(12), 3768–3790.
- Zytnick, J. (2025). Imputing proxy advisor recommendations. *Journal of Empirical Legal Studies* 22(4), 525–543.

Figure 1: McCrary Manipulation test

This figure displays the McCrary density test for the running variable (voting result, in percent) around the majority threshold. The plot shows binned counts and the fitted densities on each side of the threshold.

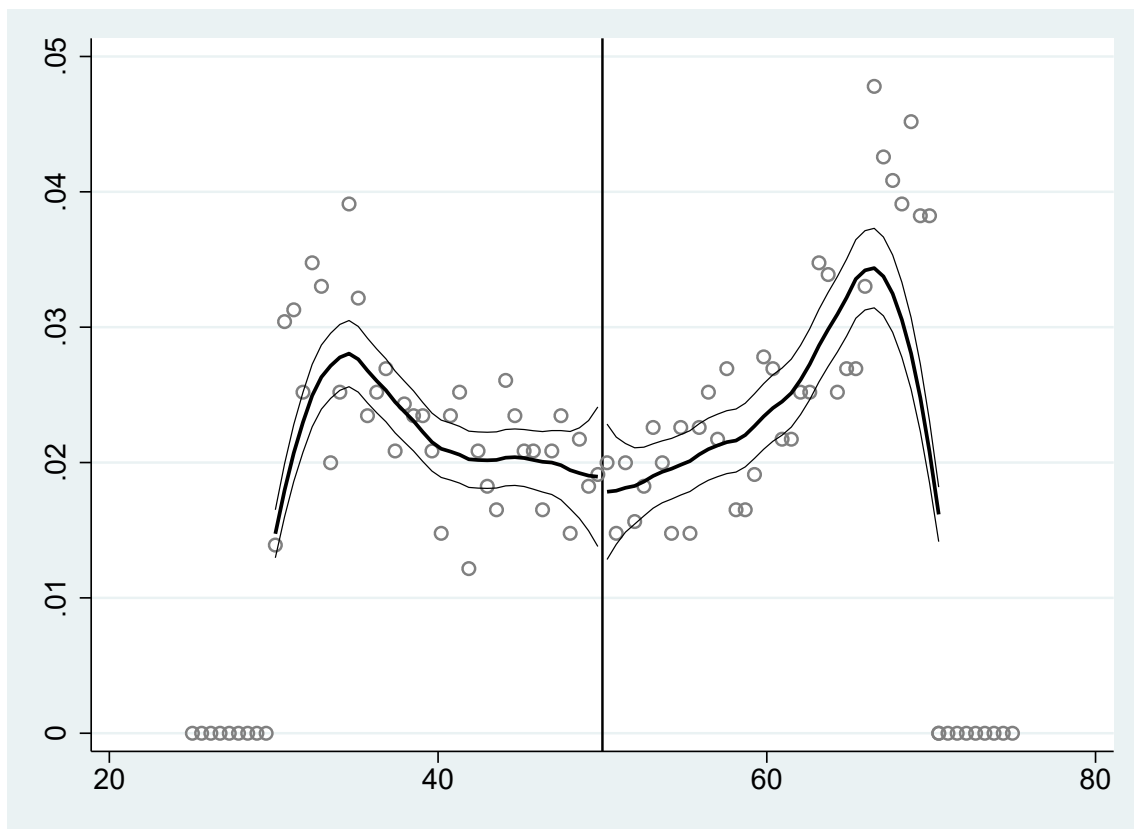


Figure 2: Long-Run CAR Patterns for AI-Aligned vs. AI-Non-Aligned Votes

This figure plots average cumulative abnormal returns from 10 trading days before the vote date to 180 trading days after the vote for two subsamples in the out-of-sample period. The sample is restricted to proposals where the absolute difference in predicted CARs between the two possible voting outcomes (for or against) is among the top 50% of differences as predicted by the AI model and which voting result is within a 2.5% margin of the majority threshold. The red line represents proposals where shareholders voted in alignment with the AI-recommended outcome (“AI-optimal”), and the blue line represents proposals where shareholders voted against the AI recommendation (“AI-suboptimal”). The trajectories illustrate differences in long-run market reactions depending on whether shareholders follow AI recommendation.

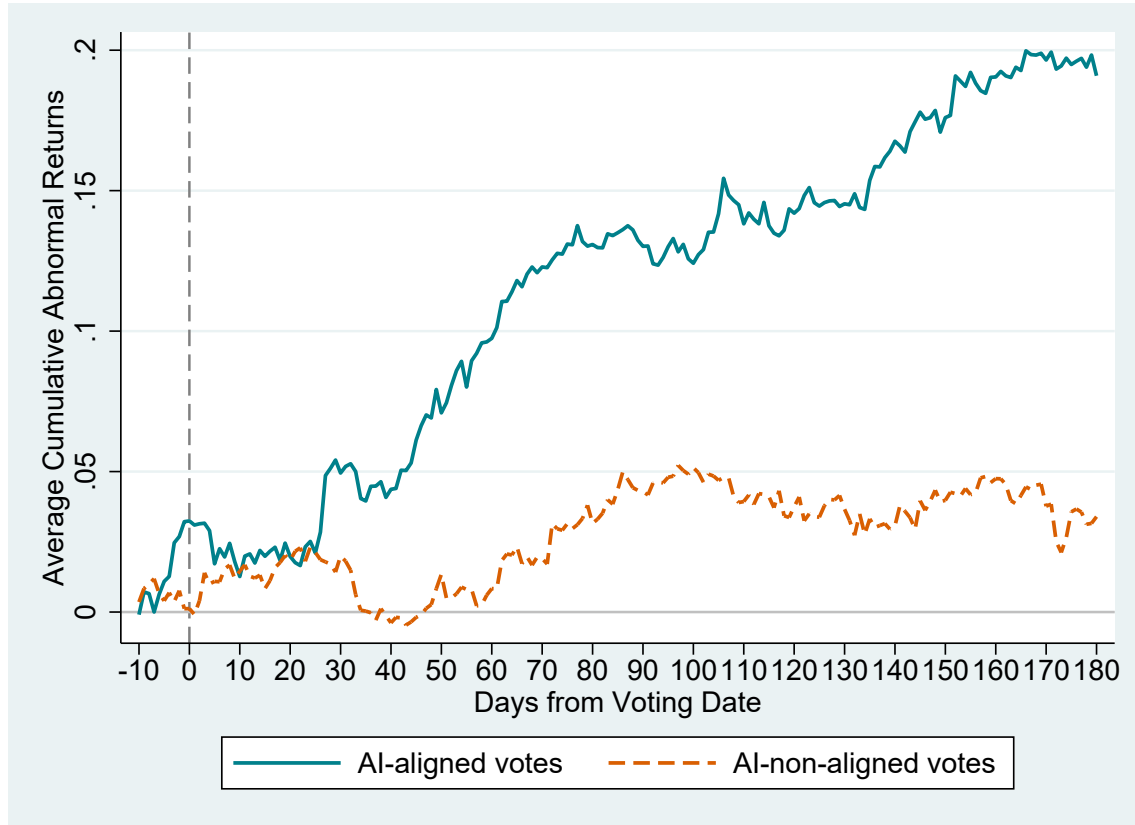


Table 1: Descriptive Statistics

This table presents descriptive statistics for proposal and firm-level variables for the set of proposals which received a shareholder support between 30% and 70%. Columns 1 through 6 report the mean, median, standard deviation, minimum, maximum, and the number of observations (N), respectively. Panel A presents proposal characteristics. Panel B presents firm characteristics. All continuous variables were winsorized at the 1% and 99%. *Total assets*, *Sales*, *Labor productivity*, and *Capital expenditures* are expressed in billions of USD. Variable definitions are provided in Table A1.

	Mean (1)	Median (2)	Std. Dev. (3)	Maximum (4)	Minimum (5)	N (6)
Panel A: Proposal characteristics						
Pass	0.59	1.00	0.49	0.00	1.00	4270
Votes ratio	0.52	0.54	0.13	0.30	0.70	4270
Shareholder proposal	0.41	0.00	0.49	0.00	1.00	4270
Management recommendation	0.59	1.00	0.49	0.00	1.00	4252
ISS recommendation	0.48	0.00	0.50	0.00	1.00	4270
GL recommendation	0.57	1.00	0.50	0.00	1.00	4270
Articles & Bylaws proposal	0.11	0.00	0.31	0.00	1.00	4270
Board & Governance proposal	0.41	0.00	0.49	0.00	1.00	4270
Capital Structure & M&A proposal	0.01	0.00	0.11	0.00	1.00	4270
Audit & Financial Policy proposal	0.00	0.00	0.03	0.00	1.00	4270
Compensation proposal	0.24	0.00	0.43	0.00	1.00	4270
ESG proposal	0.11	0.00	0.31	0.00	1.00	4270
Other proposal	0.13	0.00	0.33	0.00	1.00	4270
Panel B: Firm characteristics						
Total assets	45.88	8.75	110.05	0.24	689.92	2837
Sales	19.18	4.31	36.68	0.13	179.04	2837
Book-to-market ratio	1.97	1.04	2.75	0.12	15.66	2837
Institutional ownership (%)	0.81	0.83	0.16	0.02	1.14	2833
Leverage	0.32	0.31	0.21	0.00	0.96	2837
ROA	0.05	0.04	0.08	-0.22	0.28	2837
ROE	0.42	0.27	0.63	-0.21	4.30	2284
Profit margin	0.08	0.08	0.17	-0.66	0.54	2837
Labor productivity	0.94	0.45	1.60	0.05	10.83	2800
Capital expenditures	1.04	0.13	2.43	0.00	13.79	2832
Dividend yield	0.02	0.02	0.02	0.00	0.11	2837
Tobin's Q	2.30	1.72	1.67	0.82	9.60	2547
One-year return	0.07	0.04	0.36	-0.67	1.41	2837
Governance pillar score	5.12	5.20	1.27	2.00	7.90	2837
Environmental pillar score	5.17	5.10	2.39	0.00	10.00	2837
Social pillar score	4.31	4.30	1.48	0.50	8.20	2837
Analysts' recommendation	2.34	2.31	0.44	1.33	3.50	2837

Table 2: Out-of-sample CARs - AI recommendations

This table reports average cumulative abnormal returns (CARs) over a three-day window surrounding the vote date. “AI-aligned votes” refers to the subsample where shareholders voted in alignment with the model-recommended outcome (i.e., the outcome with the higher predicted CAR). “AI-non-aligned votes” refers to the subsample where shareholders voted against the model-recommended outcome. “Difference” is the spread in predicted CARs between the two outcomes and tests whether the market reaction for the AI-aligned sample is larger than for the AI-non-aligned sample. Columns 1–4 report results for proposals with shareholder support within 10%, 7.5%, 5%, and 2.5% margins around the majority threshold, respectively. Column 5 reports results for proposals within the optimal bandwidth of 3.8%, obtained following `calonico2020optimal`. Column 6 reports results for the full sample of proposals within a 20% margin around the majority threshold and includes control variables as well as year and proposal-topic fixed effects. Standard errors, reported in brackets, are robust to heteroskedasticity and clustered at the firm level. One, two, and three asterisks denote statistical significance at the 1%, 5%, and 10% levels, respectively.

	+/-10%	+/-7.5%	+/-5%	+/-2.5%	Optimal bandwidth	Full model
	(1)	(2)	(3)	(4)	(5)	(6)
AI-aligned votes	0.0019 (0.0023)	0.0015 (0.0025)	0.0003 (0.0030)	0.0066 (0.0040)	0.0023 (0.0030)	-0.0011 (0.0018)
AI-non-aligned votes	-0.0056** (0.0026)	-0.0078*** (0.0028)	-0.0094*** (0.0030)	-0.0086* (0.0046)	-0.0094*** (0.0035)	-0.0043* (0.0022)
Difference	0.0076*** (0.0032)	0.0093*** (0.0035)	0.0097*** (0.0040)	0.0151*** (0.0062)	0.0117*** (0.0044)	0.0047* (0.0031)
Observations	816	615	388	181	296	2053
Controls	No	No	No	No	No	Yes
Year fixed effects	No	No	No	No	No	Yes
Topic fixed effects	No	No	No	No	No	Yes

Table 3: Out-of-sample CARs - AI vs. Proxy advisors

This table reports average cumulative abnormal returns (CARs) over a three-day window surrounding the vote date. The analysis compares shareholder voting alignment with proxy advisor recommendations—ISS and Glass Lewis—and AI recommendations. “ISS-aligned” and “GL-aligned” refer to subsamples where shareholders voted in alignment with the respective proxy advisor’s recommendation, while “ISS-non-aligned” and “GL-non-aligned” refer to subsamples where shareholders voted against the recommendation. “AI-aligned” refers to the subsample where shareholders voted in alignment with the AI-recommended outcome (i.e., the outcome associated with the higher predicted CAR). Columns labelled “Difference” measures the spread in CARs between aligned and non-aligned votes for each proxy advisor and provided one-sided tests of whether the market reaction is larger when shareholders follow the recommendation. Additional regressions include both AI-aligned and proxy advisor-aligned indicators to assess whether the market reaction is stronger when shareholders follow AI-recommended outcomes relative to proxy advisor recommendations. The rows labelled “Difference AI – ISS” and “Difference AI – GL” report the estimated difference in CARs between AI-aligned and ISS-aligned votes and between AI-optimal and GL-aligned votes, respectively. These differences are computed as linear combinations of coefficients from the joint regressions. A positive and statistically significant value indicates that following AI’s recommendation yields higher CARs than following the proxy advisor’s recommendation within the specified bandwidth. Columns 1–4 report results for proposals with shareholder support within 10%, 7.5%, 5%, and 2.5% margins around the majority threshold, respectively. Column 5 reports results for proposals within the optimal bandwidth of 3.8%, obtained following *calonico2020optimal*. Column 6 reports results for the full sample of proposals within a 20% margin around the majority threshold and includes control variables as well as year fixed effects. Standard errors, reported in brackets, are robust to heteroskedasticity and clustered at the firm level. One, two, and three asterisks denote statistical significance at the 1%, 5%, and 10% levels, respectively.

(a) Panel A: AI vs. ISS

	+/-10%	+/-7.5%	+/-5%	+/-2.5%	Optimal band- width	Full model
	(1)	(2)	(3)	(4)	(5)	(6)
Panel A: AI vs. ISS						
ISS-aligned votes	-0.0037 (0.0031)	-0.0047 (0.0037)	-0.0053* (0.0031)	-0.0065 (0.0046)	-0.0058* (0.0033)	-0.0021 (0.0021)
ISS-non-aligned votes	-0.0009 (0.0021)	-0.0025 (0.0022)	-0.0039 (0.0028)	0.0024 (0.0040)	-0.0018 (0.0031)	-0.0027 (0.0018)
Difference	-0.0028 (0.0035)	-0.0022 (0.0040)	-0.0014 (0.0040)	-0.0090* (0.0062)	-0.0040 (0.0044)	0.0012 (0.0025)
AI-aligned	0.0074** (0.0034)	0.0093*** (0.0036)	0.0099** (0.0042)	0.0139** (0.0064)	0.0114** (0.0047)	0.0050 (0.0031)
ISS-aligned	-0.0009 (0.0036)	-0.0000 (0.0041)	0.0008 (0.0042)	-0.0059 (0.0063)	-0.0014 (0.0047)	0.0019 (0.0026)
Difference	0.0083** (0.0043)	0.0094** (0.0049)	0.0091** (0.0049)	0.0199*** (0.0076)	0.0129*** (0.0052)	0.0031 (0.0037)
Observations	816	615	388	181	296	2053
Controls	No	No	No	No	No	Yes
Year fixed effects	No	No	No	No	No	Yes
Topic fixed effects	No	No	No	No	No	Yes

Table 3: Out-of-sample CARs - AI vs. Proxy advisors (Continued)

(b) Panel B: AI vs. Glass Lewis

	+/-10%	+/-7.5%	+/-5%	+/-2.5%	Optimal band- width	Full model
	(1)	(2)	(3)	(4)	(5)	(6)
GL-aligned votes	-0.0003 (0.0033)	-0.0028 (0.0036)	-0.0052 (0.0035)	-0.0028 (0.0052)	-0.0047 (0.0040)	-0.0026 (0.0020)
GL-non-aligned votes	-0.0027 (0.0018)	-0.0034* (0.0020)	-0.0038 (0.0026)	0.0003 (0.0037)	-0.0024 (0.0029)	-0.0025 (0.0018)
Difference	0.0025 (0.0036)	0.0006 (0.0039)	-0.0014 (0.0040)	-0.0031 (0.0064)	-0.0023 (0.0049)	-0.0005 (0.0023)
AI-aligned	0.0081** (0.0032)	0.0096*** (0.0035)	0.0098** (0.0042)	0.0151** (0.0065)	0.0117*** (0.0045)	0.0047 (0.0031)
GL-aligned	0.0037 (0.0035)	0.0021 (0.0038)	0.0004 (0.0042)	-0.0002 (0.0067)	-0.0001 (0.0051)	-0.0003 (0.0023)
Difference	0.0044 (0.0047)	0.0076* (0.0050)	0.0094** (0.0050)	0.0153** (0.0076)	0.0118** (0.0059)	0.0050* (0.0038)
Observations	816	615	388	181	296	2053
Controls	No	No	No	No	No	Yes
Year fixed effects	No	No	No	No	No	Yes
Topic fixed effects	No	No	No	No	No	Yes

Table 4: Out-of-sample CARs - AI vs. AI+PA recommendations

This table studies whether providing proxy advisors' recommendations to the AI model allows the model to make better recommendations. It reports average cumulative abnormal returns (CARs) over a three-day window surrounding the vote date. "AI-aligned" refers to the subsample where shareholders voted in alignment with the benchmark model recommendation. In Panel A, "AI+ISS-aligned" refers to the subsample where shareholders voted in alignment with the model-recommended outcome, where the model includes ISS recommendation. "Difference" is the spread in predicted CARs between the two outcomes and provides a one-sided test of whether the market reaction for the AI-aligned sample is larger than for the AI+ISS-aligned sample. In Panel B, "AI+GL-aligned" refers to the subsample where shareholders voted in alignment with the model-recommended outcome, where the model includes Glass Lewis recommendation. "Difference" is the spread in predicted CARs between the two outcomes and tests whether the market reaction for the AI-aligned sample is larger than for the AI+GL-aligned sample. Columns 1–4 report results for proposals with shareholder support within 10%, 7.5%, 5%, and 2.5% margins around the majority threshold, respectively. Column 5 reports results for proposals within the optimal bandwidth of 3.8%, obtained following calonico2020optimal. Column 6 reports results for the full sample of proposals within a 20% margin around the majority threshold and includes control variables as well as year and proposal-topic fixed effects. Standard errors, reported in brackets, are robust to heteroskedasticity and clustered at the firm level. One, two, and three asterisks denote statistical significance at the 1%, 5%, and 10% levels, respectively.

	+/-10%	+/-7.5%	+/-5%	+/-2.5%	Optimal bandwidth	Full model
	(1)	(2)	(3)	(4)	(5)	(6)
Panel A: AI vs. AI+ISS						
AI-aligned	0.0065*	0.0083**	0.0081*	0.0091	0.0089*	0.0034
	(0.0036)	(0.0037)	(0.0045)	(0.0075)	(0.0048)	(0.0036)
AI+ISS-aligned	0.0024	0.0023	0.0036	0.0142*	0.0067	0.0039
	(0.0037)	(0.0038)	(0.0042)	(0.0076)	(0.0048)	(0.0036)
Difference	0.0042	0.0060	0.0045	-0.0052	0.0022	-0.0005
	(0.0063)	(0.0061)	(0.0073)	(0.0136)	(0.0081)	(0.0063)
Panel B: AI vs. AI+GL						
AI-aligned	0.0077**	0.0095***	0.0101**	0.0171***	0.0121***	0.0045
	(0.0033)	(0.0036)	(0.0041)	(0.0062)	(0.0044)	(0.0033)
AI+GL-aligned	-0.0013	-0.0012	-0.0036	-0.0096	-0.0045	0.0018
	(0.0034)	(0.0036)	(0.0038)	(0.0062)	(0.0042)	(0.0027)
Difference	0.0090**	0.0107**	0.0136**	0.0267***	0.0167***	0.0027
	(0.0052)	(0.0055)	(0.0063)	(0.0097)	(0.0065)	(0.0051)
Observations	816	615	388	181	296	2053
Controls	No	No	No	No	No	Yes
Year fixed effects	No	No	No	No	No	Yes
Topic fixed effects	No	No	No	No	No	Yes

Table 5: Support rates across shareholders, AI, and proxy advisors

This table presents support rates for all proposals, shareholder-sponsored proposals, and management-sponsored proposals. Column 1 reports the percentage of proposals that shareholders supported, Column 2 reports the percentage of proposals AI recommended support, Column 3 reports the percentage of proposals ISS recommended support, and Column 4 reports the percentage of proposals Glass Lewis recommended support. The statistics are computed for the out-of-sample period.

	Shareholders (1)	AI (2)	ISS (3)	GL (4)
All	0.56	0.22	0.51	0.58
Shareholder proposals	0.16	0.12	0.91	0.69
Management proposals	0.87	0.29	0.20	0.49

Table 6: Determinants of proposal support: Shareholders, AI, and proxy advisors

This table reports results from regressions examining the determinants of proposal support across four decision-makers. In Column 1, the dependent variable is an indicator equal to 1 if shareholders support the proposal. In Column 2, the dependent variable is an indicator equal to 1 if the AI model recommends supporting the proposal. In Column 3, the dependent variable is an indicator equal to 1 if ISS recommends supporting the proposal. In Column 4, the dependent variable is an indicator equal to 1 if Glass Lewis recommends supporting the proposal. Year and proposal-topic fixed effects are included in all specifications. Standard errors, reported in brackets, are robust to heteroskedasticity and clustered at the firm level. The regression are estimated on the out-of-sample period. One, two, and three asterisks denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	Shareholder support (1)	AI support (2)	ISS support (3)	GL support (4)
Proposal characteristics:				
Shareholder proposal	-0.6689*** (0.0364)	-0.0405 (0.0334)	0.5643*** (0.0427)	0.3258*** (0.0538)
Board & Governance	0.1024** (0.0403)	0.0191 (0.0380)	-0.1822*** (0.0451)	0.3196*** (0.0659)
Capital Structure & M&A	-0.1551*** (0.0600)	-0.1468 (0.1034)	0.0177 (0.0595)	0.3777*** (0.0941)
Compensation	-0.0737* (0.0443)	-0.0205 (0.0381)	-0.3103*** (0.0451)	0.2052*** (0.0741)
ESG	0.0664 (0.0438)	-0.0662* (0.0347)	-0.0645** (0.0275)	0.1250** (0.0619)
Other	0.0275 (0.0402)	-0.0550 (0.0395)	-0.0581* (0.0350)	0.1932*** (0.0579)
Firm characteristics:				
Governance score	-0.0113 (0.0071)	-0.0246** (0.0115)	0.0053 (0.0093)	0.0195 (0.0128)
Environmental score	0.0038 (0.0037)	-0.0015 (0.0061)	-0.0023 (0.0046)	-0.0092 (0.0066)
Social score	-0.0070 (0.0057)	-0.0211** (0.0102)	0.0139* (0.0071)	-0.0074 (0.0099)
Analysts' recommendations	-0.0011 (0.0193)	-0.0045 (0.0311)	0.0143 (0.0354)	0.0018 (0.0353)
Log(Market capitalisation)	-0.0221*** (0.0056)	-0.0284*** (0.0104)	-0.0017 (0.0080)	-0.0035 (0.0104)
ROA	0.3229*** (0.1218)	-0.5279** (0.2100)	0.4360*** (0.1656)	-0.3351 (0.2223)
Book leverage	-0.0242 (0.0415)	0.1210** (0.0544)	-0.0807** (0.0344)	-0.0062 (0.0728)
Book-to-market ratio	-0.0018 (0.0014)	-0.0070** (0.0028)	0.0014 (0.0022)	0.0044 (0.0028)
Profit margin	-0.0296 (0.0420)	-0.1419** (0.0703)	-0.0929* (0.0537)	0.0951 (0.0764)
One year return	0.0120 (0.0233)	-0.0372 (0.0267)	-0.0036 (0.0207)	-0.0326 (0.0279)
Dividend yield	0.0656 (0.3404)	2.2502*** (0.6139)	-0.7231 (0.4956)	0.0832 (0.7019)
Constant	1.0861*** (0.0845)	0.8137*** (0.1403)	0.3995*** (0.1228)	0.2986* (0.1627)
Observations	2053	2053	2053	2053
Adjusted R^2	0.533	0.157	0.529	0.079
Year fixed effects	Yes	Yes	Yes	Yes
Topic fixed effects	Yes	Yes	Yes	Yes

Table 7: Alignment of shareholder votes with AI and proxy advisors' recommendations

This table presents the percentage of proposals where shareholders' votes align with recommendations from AI, ISS, and Glass Lewis. Column 1 reports the percentage of proposals where shareholders followed the AI-recommended outcome, Column 2 reports the percentage where shareholders followed ISS's recommendation, and Column 3 reports the percentage where shareholders followed Glass Lewis's recommendation. The first row shows alignment rates for all proposals, while the second and third rows report rates for shareholder-sponsored and management-sponsored proposals separately. The statistics are computed for the out-of-sample period.

	Follows AI (1)	Follows ISS (2)	Follows GL (3)
All proposals	53.4	28.1	50.3
Management proposals	34.7	31.5	56.8
Shareholder proposals	77.8	23.6	41.8

Table 8: Determinants of shareholder alignment with AI and proxy advisors' recommendations

This table reports results from regressions examining the determinants of whether shareholders' votes align with recommendations from AI, ISS, and Glass Lewis. In Column 1, the dependent variable is an indicator equal to 1 if shareholders followed the AI-recommended outcome. In Column 2, the dependent variable is an indicator equal to 1 if shareholders followed ISS's recommendation. In Column 3, the dependent variable is an indicator equal to 1 if shareholders followed Glass Lewis's recommendation. Year and proposal-topic fixed effects are included in all specifications. Standard errors, reported in brackets, are robust to heteroskedasticity and clustered at the firm level. The regression are estimated on the out-of-sample period. One, two, and three asterisks denote statistical significance at the 10%, 5%, and 1% levels, respectively.

	Follow AI (1)	Follow ISS (2)	Follow GL (3)
Proposal characteristics:			
Shareholder proposal	0.4653*** (0.0378)	0.0253 (0.0510)	-0.2568*** (0.0501)
Board & Governance	-0.0502 (0.0458)	0.2229*** (0.0563)	-0.2127*** (0.0605)
Capital Structure & M&A	0.0061 (0.0816)	0.0896 (0.0930)	-0.3479*** (0.0907)
Compensation	0.0144 (0.0512)	0.2174*** (0.0587)	-0.1567** (0.0662)
ESG	-0.0397 (0.0525)	0.1302*** (0.0483)	0.0071 (0.0591)
Other	0.0144 (0.0504)	0.0842* (0.0493)	-0.1138** (0.0561)
Firm characteristics:			
Governance score	-0.0152 (0.0118)	0.0275*** (0.0106)	0.0362*** (0.0130)
Environmental score	-0.0012 (0.0058)	-0.0009 (0.0055)	0.0097 (0.0068)
Social score	-0.0217** (0.0093)	0.0280*** (0.0082)	0.0151* (0.0087)
Analysts' recommendations	-0.0344 (0.0310)	0.0505 (0.0337)	-0.0002 (0.0326)
Log(Market capitalisation)	-0.0082 (0.0098)	-0.0068 (0.0096)	-0.0164 (0.0114)
ROA	-0.6392*** (0.1842)	0.2162 (0.1948)	-0.0394 (0.2300)
Book leverage	0.0260 (0.0481)	-0.1196** (0.0479)	0.0664 (0.0560)
Book-to-market ratio	-0.0041 (0.0028)	-0.0048** (0.0021)	0.0017 (0.0029)
Profit margin	-0.0074 (0.0554)	-0.0716 (0.0602)	-0.0041 (0.0683)
One year return	-0.0332 (0.0318)	-0.0381 (0.0243)	-0.0329 (0.0288)
Dividend yield	2.1201*** (0.5954)	-0.8687* (0.5132)	0.6779 (0.6462)
Constant	0.6076*** (0.1373)	-0.1262 (0.1365)	0.5867*** (0.1604)
Observations	2053	2053	2053
Adjusted R^2	0.215	0.044	0.057
Year fixed effects	Yes	Yes	Yes
Topic fixed effects	Yes	Yes	Yes

Table 9: Counterfactual Voting Outcomes if Shareholders Followed AI Recommendations

This table reports how shareholder voting outcomes would change if investors fully followed the AI's recommendations. For each proposal topic, Columns 1 and 2 show the percentage of proposals where shareholders voted in the same direction as the AI recommendation (Against or For, respectively). Columns 3 and 4 report the percentage of proposals where shareholders voted against the AI recommendation when it advised Against or For, respectively. Column 5 presents the difference between Columns 3 and 4, capturing the net shift in voting outcomes implied by full compliance with the AI recommendation. All values are expressed in percent. Topics are sorted by their frequency in the sample.

AI recommendation	Follows AI rec.		Does not follow AI rec.		
	Against (1)	For (2)	Against (3)	For (4)	Difference (5)=(3)-(4)
CEO-Chairman Separation	80.87	0	2.61	16.52	13.91
Director Election	1.02	30.81	66.86	1.31	-65.55
Equity Compensation Plans	5.43	28.26	66.3	0	-66.3
GHG Emissions and Climate Reporting	67.92	1.89	30.19	0	-30.19
Golden Parachutes	63.64	6.82	20.45	9.09	-11.36
Labor Rights	70.64	5.5	17.43	6.42	-11.01
Political Disclosure	75	1.72	12.93	10.34	-2.59
Proxy Access	82.86	0	5.71	11.43	5.71
Say On Pay	22.59	16.57	52.41	8.43	-43.98
Special Meetings	81.37	0	13.66	4.97	-8.7
Written Consent	76.42	2.83	6.6	14.15	7.55

Table 10: Long-term impacts of following AI recommendations

This table examines the impact of following the AI's voting recommendation on *growth rates* of ROA, profit margin, ROE, Tobin's Q, sales, labor productivity, and capex. We estimate a regression discontinuity design using proposals within a bandwidth of $\pm 7.5\%$ around the majority threshold. The dependent variables in all columns are year-over-year growth measures, winsorized at the 1st and 99th percentiles by fiscal year. Panel A reports results for the fiscal year of the vote (t); Panel B reports results for the subsequent fiscal year ($t + 1$). All specifications include firm and year fixed effects, and standard errors are clustered at the firm level. Variable definitions are provided in Table A1.

	ROA	Profit margin	ROE	Tobin's Q	Sales	Labor productivity	Capex
	(1)	(2)	(3)	(4)	(5)	(6)	(7)
Panel A: Year of the vote (t)							
AI-aligned votes	-0.0019 (0.0092)	-0.0155 (0.0270)	0.1200 (0.1278)	0.0249 (0.0352)	-0.0305 (0.0430)	-0.0317 (0.0575)	-0.1002 (0.1128)
Observations	388	386	325	362	386	384	354
Year fixed effects	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Firm fixed effects	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Panel B: Year after the vote ($t + 1$)							
AI-aligned votes	0.0121** (0.0063)	0.0358*** (0.0145)	0.0717 (0.0643)	0.0088 (0.0363)	0.0586** (0.0270)	0.0607** (0.0287)	0.0643 (0.0659)
Observations	386	386	317	355	386	384	352
Year fixed effects	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Firm fixed effects	Yes	Yes	Yes	Yes	Yes	Yes	Yes

Table 11: Robustness to Table 2 - Predicted be within 20% margin

This table presents a robustness test to the analyses presented in Table 2 by using a sample of proposals that are predicted to be within the 20% margin of the majority threshold. As for Table 2, this table reports average cumulative abnormal returns (CARs) over a three-day window surrounding the vote date. “AI-aligned” refers to the subsample where shareholders voted in alignment with the model-recommended outcome (i.e., the outcome with the higher predicted CAR). “AI-non-aligned” refers to the subsample where shareholders voted against the model-recommended outcome. “Difference” is the spread in predicted CARs between the two outcomes and tests whether the market reaction for the AI-aligned sample is larger than for the AI-aligned sample. Columns 1–4 report results for proposals with shareholder support within 10%, 7.5%, 5%, and 2.5% margins around the majority threshold, respectively. Column 5 reports results for proposals within the optimal bandwidth of 3.8%, obtained following `calonico2020optimal`. Column 6 reports results for the full sample of proposals within a 20% margin around the majority threshold and includes control variables as well as year and proposal-topic fixed effects. Standard errors, reported in brackets, are robust to heteroskedasticity and clustered at the firm level. One, two, and three asterisks denote statistical significance at the 1%, 5%, and 10% levels, respectively.

	+/-10%	+/-7.5%	+/-5%	+/-2.5%	Optimal bandwidth	Full model
	(1)	(2)	(3)	(4)	(5)	(6)
AI-aligned votes	-0.0001 (0.0022)	0.0000 (0.0022)	-0.0011 (0.0025)	0.0054 (0.0036)	-0.0005 (0.0027)	0.0001 (0.0018)
AI-non-aligned votes	-0.0057** (0.0024)	-0.0070** (0.0028)	-0.0075** (0.0035)	-0.0085 (0.0055)	-0.0065 (0.0041)	-0.0034 (0.0036)
Difference	0.0056** (0.0032)	0.0070** (0.0036)	0.0064* (0.0043)	0.0140** (0.0066)	0.0071*** (0.0030)	0.0040 (0.0043)
Observations	519	398	258	120	648	999
Controls	No	No	No	No	No	Yes
Year fixed effects	No	No	No	No	No	Yes
Topic fixed effects	No	No	No	No	No	Yes

A Appendix

Table A1: Variable Definitions

The table contains the definitions and data sources of the key variables used in the paper.

Variable	Definition	Source
Proposal characteristics		
Pass	Indicator equal to 1 if the proposal receives at least 50% of votes cast, and 0 otherwise.	ISS Voting Analytics
Votes ratio	Number of votes cast in favor of the proposal over the number of votes cast.	ISS Voting Analytics
Shareholder proposal	Indicator equal to 1 if the proposal's sponsor is a shareholder, and 0 if it is management	ISS Voting Analytics
Management recommendation	Indicator equal to 1 if management recommends to vote in favor of the proposal, and 0 otherwise	ISS Voting Analytics
ISS recommendation	Indicator equal to 1 if ISS recommends to vote in favor of the proposal, and 0 otherwise	Zytnick (2025)
GL recommendation	Indicator equal to 1 if GL recommends to vote in favor of the proposal, and 0 otherwise	Zytnick (2025)
Articles & Bylaws proposal	Indicator equal to 1 if the proposal addresses changes in articles and bylaws,	Constructed using ISS Voting Analytics
Board & Governance proposal	Indicator equal to 1 if the proposal addresses a board or governance related topic,	Constructed using ISS Voting Analytics
Capital structure & M&A proposal	Indicator equal to 1 if the proposal addresses a capital structure or merger and acquisition related topic,	Constructed using ISS Voting Analytics
Audit & Financial Policy proposal	Indicator equal to 1 if the proposal addresses an audit or financial policy related topic,	Constructed using ISS Voting Analytics
Compensation proposal	Indicator equal to 1 if the proposal addresses a compensation related topic,	Constructed using ISS Voting Analytics
ESG proposal	Indicator equal to 1 if the proposal addresses an ESG related topic,	Constructed using ISS Voting Analytics
Other proposal	Indicator equal to 1 if the proposal addresses a topic that does not belong to the above categories,	Constructed using ISS Voting Analytics
Proposal characteristics		
Total assets	Total assets (in billions of USD), at .	Compustat/CRSP
Sales	Sales (in billions of USD), $sale$.	Compustat/CRSP
Book-to-market ratio	Book value of equity divided by market value of equity, $at/(prcc.f \cdot csho)$.	Compustat/CRSP
Institutional ownership (%)	Percentage of shares held by institutional investors.	13F (Thomson/Refinitiv)
Leverage	Total debt divided by total assets, $(dltt + dlc)/at$.	Compustat/CRSP
Return on assets (ROA)	Net income divided by total assets, ni/at .	Compustat/CRSP
Return on equity (ROE)	Net income divided by book equity, $ni/(ceq + txditc)$.	Compustat/CRSP
Profit margin	Net income divided by total revenues, $ni/sale$.	Compustat/CRSP
Labor productivity	Sales divided by number of employees, $sale/emp$.	Compustat/CRSP
Capital expenditures	Capital expenditures during the fiscal year (in billions of USD), $capx$.	Compustat/CRSP
Dividend yield	Dividends per share divided by lagged stock price, $(dv/csho)/prcc.f_{t-1}$.	Compustat/CRSP
Tobin's Q	Ratio of market value of assets to book value of assets, $(at + csho \cdot prcc.f - ceq + txditc)/at$.	Compustat/CRSP
One-year stock return	Buy-and-hold stock return over one year, $prcc.f_t/prcc.f_{t-1} - 1$.	Compustat/CRSP
Governance pillar score	Governance pillar score.	MSCI ESG Ratings
Environmental pillar score	Environmental pillar score.	MSCI ESG Ratings
Social pillar score	Social pillar score.	MSCI ESG Ratings
Analysts' recommendation	Mean analysts recommendation	IBES

Internet Appendix (IA)

IA1 Additional results

Table IA1: Out-of-Sample Classification Performance: Contentiousness and Voting Outcomes

This table reports the out-of-sample classification performance of the model along two dimensions. Panel A evaluates the ability of the model to identify contentious shareholder proposals, defined as proposals receiving a vote share between 30% and 70%. Panel B evaluates the ability of the model to predict the voting outcome (pass versus fail) among predicted non-contentious proposals, defined as proposals with vote shares below 30% or above 70%. In each panel, the confusion matrix compares predicted outcomes (columns) with realized outcomes (rows). Each cell reports the number of observations. Row percentages report the fraction of observations within each realized outcome category that are assigned to each predicted category. Column percentages report the fraction of observations within each predicted category that correspond to each realized outcome. The sample includes all proposals for which both realized and predicted vote shares are available.

Panel A: Prediction of Contentious Proposals

	Predicted non- contentious	Predicted contentious
Realized non-contentious	71584	710
Row %	99	1
Col %	98.4	37.6
Realized contentious	1198	1178
Row %	50.4	49.6
Column %	1.6	62.4

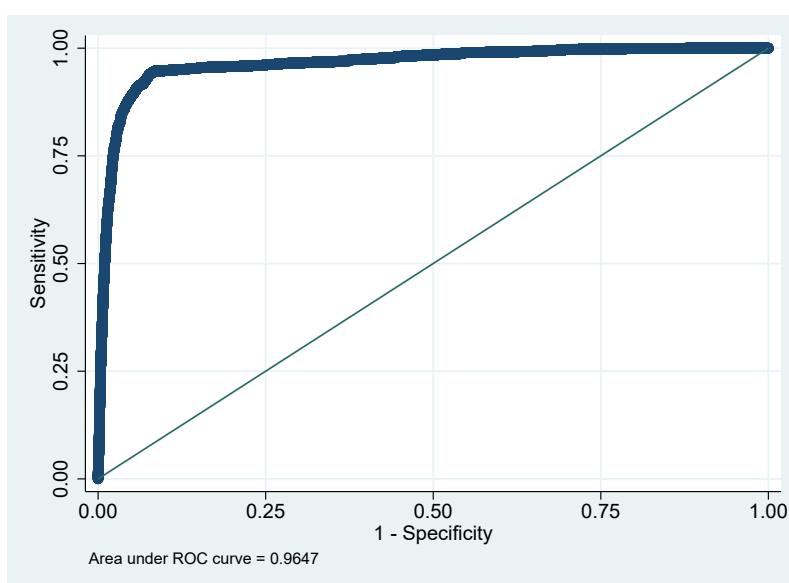
Panel B: Prediction of Voting Outcomes among Non-Contentious Proposals

	Predicted fail	Predicted pass
Realized failure	1551	99
Row %	94	6
Col %	97.4	.1
Realized success	41	71091
Row %	.1	99.9
Column %	2.6	99.9

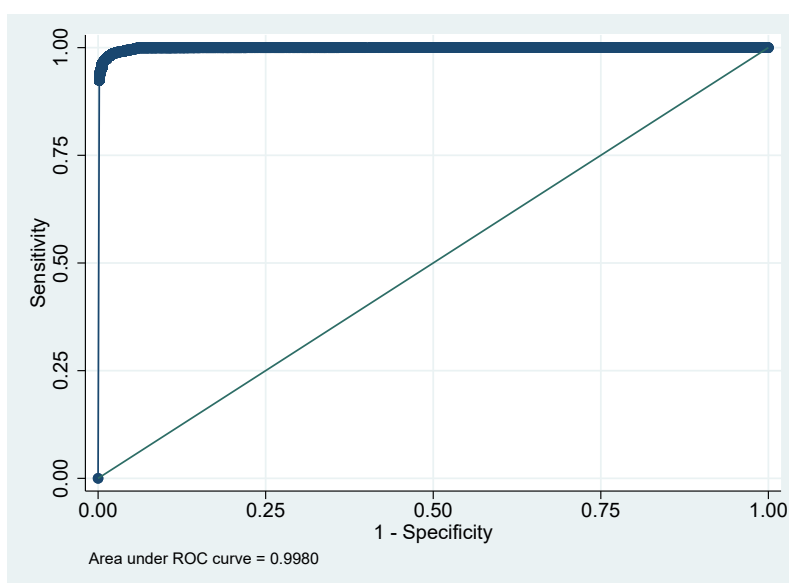
Figure IA1: Out-of-Sample ROC Curves: Contentiousness and Voting Outcomes

This figure reports the out-of-sample receiver operating characteristic (ROC) curves evaluating the classification performance of the model along two dimensions. Panel A plots the ROC curve for predicting whether a shareholder proposal is contentious, defined as receiving a vote share between 30% and 70%. Panel B plots the ROC curve for predicting the voting outcome (pass versus fail) among predicted non-contentious proposals, defined as those with vote shares below 30% or above 70%. The ROC curve plots the true positive rate against the false positive rate at different classification thresholds. The area under the curve (AUC) summarizes overall classification performance, with higher values indicating better predictive accuracy. An AUC of 0.5 corresponds to random classification, while an AUC of 1 indicates perfect prediction.

Panel A: Predicting Contentious Proposals



Panel B: Predicting Voting Outcomes among Non-Contentious Proposals



IA2 Robustness tests

Table IA2: Robustness to Table 2 - Excluding cases with confounding proposals

This table presents a robustness test to the analyses presented in Table 2 by limiting the sample to cases where there is either one proposal with a result within a 20% margin of the majority threshold or when if the AI recommendations were followed for all the proposals or for none of them. As for Table 2, this table reports average cumulative abnormal returns (CARs) over a three-day window surrounding the vote date. “AI-aligned” refers to the subsample where shareholders voted in alignment with the model-recommended outcome (i.e., the outcome with the higher predicted CAR). “AI-non-aligned” refers to the subsample where shareholders voted against the model-recommended outcome. “Difference” is the spread in predicted CARs between the two outcomes and tests whether the market reaction for the AI-aligned sample is larger than for the AI-non-aligned sample. Columns 1–4 report results for proposals with shareholder support within 10%, 7.5%, 5%, and 2.5% margins around the majority threshold, respectively. Column 5 reports results for proposals within the optimal bandwidth of 3.8%, obtained following `calonico2020optimal`. Column 6 reports results for the full sample of proposals within a 20% margin around the majority threshold and includes control variables as well as year and proposal-topic fixed effects. Standard errors, reported in brackets, are robust to heteroskedasticity and clustered at the firm level. One, two, and three asterisks denote statistical significance at the 1%, 5%, and 10% levels, respectively.

	+/-10%	+/-7.5%	+/-5%	+/-2.5%	Optimal bandwidth	Full model
	(1)	(2)	(3)	(4)	(5)	(6)
AI-aligned votes	0.0047* (0.0025)	0.0039 (0.0027)	0.0028 (0.0031)	0.0081* (0.0043)	0.0051 (0.0032)	0.0004 (0.0019)
AI-non-aligned votes	-0.0057* (0.0031)	-0.0073** (0.0033)	-0.0084** (0.0034)	-0.0093* (0.0053)	-0.0081** (0.0040)	-0.0034 (0.0027)
Difference	0.0104*** (0.0040)	0.0112*** (0.0043)	0.0112*** (0.0046)	0.0175*** (0.0069)	0.0132*** (0.0051)	0.0058* (0.0040)
Observations	616	462	286	128	210	1579
Controls	No	No	No	No	No	Yes
Year fixed effects	No	No	No	No	No	Yes
Topic fixed effects	No	No	No	No	No	Yes