

Augmenting the Operations Manager with a Prediction Machine

Tao Lu¹ and Brian Tomlin²

¹School of Business, University of Connecticut, Storrs, CT 06269

²Tuck School of Business, Dartmouth College, Hanover, NH 03755

June 4, 2025 (latest version April 23, 2026)

Abstract

Demand forecasting is central to operations management, and firms look to Artificial Intelligence (AI) enabled forecasting engines (“prediction machines”) to augment their managers’ own forecasting capabilities and thus improve sales-and-operations planning outcomes. Deployment of a prediction machine may cause an unintended reduction in a manager’s own forecasting effort which in turn diminishes the value of machine adoption. We model a firm facing uncertain demand that delegates a procurement quantity decision to a human manager who can exert effort to generate a demand prediction. The firm deploys a machine that provides the manager with a demand-prediction signal. Adopting a Bayesian persuasion approach, we show that partially disclosing the machine’s prediction – either downplaying high predictions or exaggerating low predictions – can be optimal, depending on the product’s cost-to-revenue ratio. Unless the machine is significantly more accurate than the human, it can be optimal to disclose less information as the machine’s prediction capability increases; however, the resulting market-state informativeness does not decline as the machine capability increases. As an alternative to information design, we explore a contracting approach that optimizes the manager’s compensation parameters while fully disclosing machine information. We establish that an information approach outperforms a contracting approach unless managerial effort cost is low.

Key words: Information Design; Artificial Intelligence; Demand Forecasting; Newsvendor

1 Introduction

Demand forecasting is fundamental to business operations, and firms are increasingly looking at AI to improve their demand predictions. In some cases, AI is used to replace the human in forecasting

but oftentimes it is used to augment the human (Revilla et al., 2023). The successful deployment of an AI forecasting engine (“prediction machine”) typically depends on a human manager being in the loop to refine the machine’s prediction (Nair and Saenz, 2024). This reflects the fact that managers often have uncoded market insights that are not captured by the machine (Ibrahim et al., 2021; Siems and Spiliotopoulou, 2024).

A perspective that “the best forecasts often come not from replacing humans with algorithms, but from combining them” (Ibrahim et al., 2021, p2314) requires firms to consider the impact on behavior when using prediction machines to augment rather than replace humans. Concerns have been raised across a wide array of contexts (e.g., customer support, education, recruitment) that AI adoption might lead to a reduction in human effort (Ni et al., 2024; Ahmad et al., 2023; Dell’Acqua, 2022). There is some evidence that human decision makers tend to overrely on AI machine advice and that this overreliance leads to lower cognitive engagement and inefficient performance (Alexander et al., 2018; Klingbeil et al., 2024). AI adoption does not always lead to effort reduction (Ni et al., 2024) and the effect on effort may depend on the quality of an AI prediction (Dell’Acqua, 2022).

An operations manager making a procurement quantity decision in the face of demand uncertainty may exert less of their own forecasting effort if provided with a prediction machine. The firm needs to somehow account for (and possibly counteract) the effort-dampening potential of AI if it wants to take full advantage of the forecasting potential of AI. That is the problem we study in this paper. In particular, we consider a firm facing uncertain demand that delegates a purchase quantity decision to a human manager who can exert effort to generate a demand prediction. To improve its monetary profit, the firm deploys a machine that provides the manager with a signal of its demand prediction. Reflecting a common practice in which AI is used to augment but not replace human decision-making, the manager is not subservient to AI in our setting. The manager chooses whether or not to exert effort to also generate their own prediction. Based on the machine signal, or on a combination of machine signal and human predictions (if effort is exerted), the manager determines a procurement quantity.

We first consider two benchmark settings: no machine and a full-disclosure machine that communicates its actual prediction. We characterize the effort-exerting conditions with and without a machine, thereby identifying the conditions under which machine deployment can dampen the manager’s prediction effort and thus reduce the potential benefit of machine deployment. Then, we explore whether the firm can mitigate the effort dampening effect by tailoring (or tuning) the machine’s communication to the manager. Adopting a Bayesian persuasion approach, we characterize

the machine’s optimal signaling (messaging) policy, where a policy specifies a mapping from the realized machine prediction to a probability distribution for the machine signal communicated to the manager. A full-disclosure policy is one in which the machine always communicates its actual prediction. A partial-disclosure policy is one in which the machine probabilistically communicates something different from its actual prediction. For the reader uncomfortable with the notion that the firm would deploy a machine that might misrepresent its prediction, as we discuss later, a partial-disclosure policy is equivalent to the firm adjusting the machine’s true positive and true negative probabilities and implementing full disclosure (with a less-informative machine).

We establish that partial disclosure can be optimal because it can (at times) preserve managerial effort when full disclosure would dampen it. We fully characterize the firm’s optimal machine disclosure policy as a function of the product’s cost-to-revenue ratio. For high-margin products, the firm should downplay high predictions from the machine (unless managerial effort cost is low); whereas for low-margin products, the firm should exaggerate low predictions from the machine (again unless managerial effort cost is low). For medium-margin products, regardless of effort cost, the firm should neither downplay nor exaggerate the machine prediction because the information loss from obfuscation outweighs the benefit of preserving managerial effort. Interestingly, unless the machine’s prediction capability is sufficiently high, it can be optimal for the firm to disclose less information (i.e., engage in more obfuscation) as the machine prediction capability increases. An alternative approach to resolving the effort-withdrawal concern would be for the firm to optimally adjust the manager’s compensation scheme when deploying a full-disclosure prediction machine. In an extension, we identify the conditions in which our informational approach (partial disclosure) outperforms the contracting approach (compensation design). Among other extensions, we analyze a setting where the machine has a lower prediction capability than the manager and we prove that partial disclosure is optimal even for medium-margin products.

We establish in a sales-and-operations planning context that firms are right to be concerned about the potential for managerial effort reduction when AI is introduced. Importantly, our results provide actionable guidance on how to address effort withdrawal. Firms should implement an information-design based approach unless they perceive managerial effort cost to be low, in which case compensation design can outperform information design. When implementing the information-design approach, the best strategy depends on the product’s cost-to-revenue ratio. In particular, for medium ratios the firm should deploy the prediction machine (at its maximum capability) and remove the human from the forecasting loop; at very high or very low ratios the firm can deploy the

prediction machine (again at its maximum capability) and still benefit from managerial effort; but at high or low ratios the firm should only partially disclose the machine prediction by introducing bias – either exaggerating or downplaying its prediction – to retain managerial effort. In this last case, as the machine’s maximum capability improves, the firm should engage in more obfuscation (unless this maximum exceeds a threshold), but we note that the informativeness of the optimally tuned market-state prediction does not decline.

2 Literature Review

Our work adds to the emerging literature on human and machine collaboration in operations management (Dai and Swaminathan, 2025). Existing literature has explored in various context such as procurement (Cui et al., 2022; Li and Li, 2022), health care (Dai and Tayur, 2022; Dai and Singh, 2024; Hou et al., 2024; Dai and Singh, 2025), process automation (Beer et al., 2024), and enterprise generative AI design (Dai and Taylor, 2025).

More relevant to our work are those papers that explore the provision of machine-generated predictions to a human decision maker in an operations context. Through a behavioral lens, Ibrahim et al. (2021) compare two approaches to integrating human judgment with prediction algorithms, and their findings suggest that asking a human to adjust the algorithm’s prediction, rather than providing a direct forecast, leads to higher predictive accuracy. Balakrishnan et al. (2025) identify an advice-weighting behavior in which people take a weighted average of their own forecast and that of a machine, and propose ways to mitigate the resulting biases. de Véricourt and Gurkan (2023) consider a repeated decision making process and find that the human may fail to learn the machine’s true accuracy in the long run due to verification bias. Based on a rational inattention framework, Boyacı et al. (2024) study the impact of providing machine predictions to a human decision-maker who can exert cognitive effort to assess information that the machine cannot capture. They find that machine input always improves decision accuracy but may increase false positive errors. Our work complements this growing literature by asking whether we should fully disclose machine predictions to human decision makers, given that the information disclosure may affect the human’s incentive to make their own judgment.

In a concurrent paper to ours, Guan et al. (2025) also study a problem in which a principal delegates a forecasting task to a human agent and can provide a prediction machine to assist the agent. They show that the principal might reduce machine precision to incentivize the agent’s

effort. Although we and Guan et al. (2025) both examine the effect of machine predictions on human effort, the papers differ in several important ways. We embed the prediction machine in a sales-and-operations context in which an operations manager determines both forecasting effort and an order quantity. This difference has important implications because both the value of forecast information and managerial incentives depend critically on the product’s profit margin. Moreover, our framework allows the firm to control the prediction bias, downplaying or exaggerating machine predictions based on the profit margin. Additionally, our context also allows us to explore the impact of demand variance on effort and information disclosure.

We adopt an information design framework (Kamenica and Gentzkow, 2011), and as such, our work relates to the broader literature on information design in operations (Candogan, 2020). Relevant studies have considered information problems in revenue management (Drakopoulos et al., 2021), public warning policies (Alizamir et al., 2020; De Véricourt et al., 2021), platform economics (Papanastasiou et al., 2018; Bimpikis et al., 2024), smart navigation (Chen et al., 2023), and agricultural interventions (Farahani et al., 2024). The papers on the disclosure of demand information are particularly related to our work. Bimpikis and Mantegazza (2023) study a two-sided platform disclosing demand information to the supply side, thereby influencing potential suppliers’ entry and pricing decisions. Candogan and Gurkan (2024) consider a two-tier supply chain in which a retailer determines an information disclosure policy of demand information, and based on the disclosed information, a supplier chooses a production quantity before actual demand is realized. In our problem, different from Candogan and Gurkan (2024), information disclosure influences not only the order quantity of the receiver (the human manager in our setting) but also their effort in eliciting separate demand information. Moreover, the sender (the firm) maximizes a profit that includes both shortage and overage costs, whereas in Candogan and Gurkan (2024) the sender (retailer) focuses on a profit that involves only the shortage cost.

Some studies have examined information design problems in which the receiver exerts costly effort to acquire or process information. Boyaci et al. (2024) study a seller’s disclosure of product quality information, after which consumers may acquire additional information about whether the product matches their tastes. In the economics literature, Lipnowski et al. (2020) consider a receiver who must process the sender-provided information at an attention cost, and consequently the receiver’s information is bounded by the sender’s disclosure. Matyskova and Montes (2023) examine a setting where, after receiving the sender’s message and forming an interim belief, the receiver can further update its belief. In their setting, both the sender and receiver can acquire full information

about the true state. In contrast, our setting features a distinct structure in which both the sender (machine) and the receiver (human) can generate predictions, but each with only limited accuracy. As such, neither the sender nor the receiver is capable of eliciting full information alone.

Additionally, our work relates to studies on information provision or acquisition for order quantity decisions. One strand of research examines a retailer sharing private demand information with a supplier to influence its production capacity decision (Özer et al., 2011; Chu et al., 2017; Lu and Tomlin, 2024). The information is shared via cheap talk, which differs from our setting where the machine designer commits to a disclosure policy *ex ante*. Another literature stream examines the impact of information with the downstream retailer making inventory decisions (e.g., Iyer et al., 2007), and some papers consider the case where the retailer chooses not only an order quantity but also an effort level to improve demand forecasts (Taylor and Xiao, 2009; Tang and Girotra, 2017).

3 Model

We consider a human manager (“they”) choosing an order quantity Q for a single-season product that faces uncertain market demand, i.e., a classic newsvendor setting. The product is purchased at a unit cost c and sells at a unit price p , where $p > c > 0$. Demand is either a_H or a_L , where $a_H > a_L \geq 0$. Let $\theta \in \{H, L\}$ represent the uncertain market state. Let λ denote the prior probability that $\theta = H$. Two-point distributed demand allows us to derive our main results in a closed-form manner, and similar assumptions have been adopted elsewhere in the operations literature, e.g., Iyer et al. (2007). In §5.3, we show that our results extend to settings in which there is residual (unresolvable by prediction) demand uncertainty associated with each of the two demand states. We note that it is common to focus on two states in the information design literature (De Véricourt et al., 2021; Candogan and Gurkan, 2024). Before placing the order, the manager can exert effort to predict the eventual market state. The firm (that employs the manager) has decided to equip the manager with a machine that also provides a prediction of the market state.

Machine and Human Predictions. The machine generates a prediction $X_m \in \{0, 1\}$ of the market state, with $X_m = 0$ ($X_m = 1$) denoting that the machine predicts the market state to be $\theta = L$ ($\theta = H$). Interpreting the high state as positive and the low state as negative, let the “true positive” and “true negative” probabilities associated with the machine prediction be denoted as b_m and b'_m respectively; that is $b_m = \Pr(X_m = 1|\theta = H)$ and $b'_m = \Pr(X_m = 0|\theta = L)$. Separate from the machine, the human manager can also generate a market-state prediction $X_h \in \{0, 1\}$.

Analogous to the machine, we use $X_h = 0$ ($X_h = 1$) to denote a low-state (high-state) prediction, and we define $b_h = \Pr(X_h = 1|\theta = H)$ and $b'_h = \Pr(X_h = 0|\theta = L)$.

In what follows, we impose some reasonable conditions on the machine and human predictions.

Assumption 1. Both X_m and X_h are unbiased estimators of the state θ , that is, the unconditional probability of a high prediction equals the prior probability of the state being high, i.e., $\Pr(X_m = 1) = \Pr(X_h = 1) = \lambda$.

A similar assumption (and a binary forecast information structure) can be found in other demand forecasting papers in the operations and marketing literatures (Iyer et al., 2007; Jiang et al., 2016). We note that for prediction $n \in \{m, h\}$, Assumption 1 implies that $1 - b'_n = \frac{\lambda}{1-\lambda}(1 - b_n)$.¹ Without loss of generality, we focus on $b_n \in (\lambda, 1]$.²

A higher b_n signifies a more informative prediction X_n . For expositional simplicity, we will refer to b_n as the accuracy of prediction source $n \in \{m, h\}$.³ The value of b_n will depend on the nature of the product and the source of the prediction (machine or human). The more novel a product is, the harder it is (for both the machine and the human) to predict demand, and so a more novel product would be associated with lower values of b_m and b_h as compared to a less novel product. Because a prediction machine can process large amounts of data more readily than a human, in the base model we focus on the case in which $b_m > b_h$, that is, the machine is more accurate than the human in its prediction. (Because empirical evidence suggests prediction machines do not necessarily outperform humans in some settings (Abolghasemi et al., 2024), we examine the case of $b_m \leq b_h$ as a model extension in §5.2 so as to cover both cases.) For later use, we introduce the metric $\beta = \frac{1-b_h}{(1-b_m)+(1-b_h)}$ which measures the accuracy of the machine relative to the human, and note that $b_m > b_h$ if and only if (iff) $\beta > 1/2$.

There can be an overlap in the information sets used by the machine and the human to generate their predictions. For example, they may both have access to past sales data of similar products. Therefore, in general, X_m and X_h may be correlated, and so we need to consider the joint probability distribution of the machine and human predictions conditional on θ , i.e., $p_{ij|\omega} = \Pr(X_m = i, X_h = j|\theta = \omega)$ for all $i, j \in \{0, 1\}$ and $\omega \in \{L, H\}$. For tractability, we assume that the probability that

¹Because X_n is unbiased by Assumption 1, we have $\lambda = \Pr(X_n = 1) = \lambda \Pr(X_n = 1|\theta = H) + (1 - \lambda) \Pr(X_n = 1|\theta = L) = \lambda b_n + (1 - \lambda)(1 - b'_n)$ for $n \in \{m, h\}$. Thus, probabilities $1 - b_n$ and $1 - b'_n$ satisfy the relation $1 - b'_n = \frac{\lambda}{1-\lambda}(1 - b_n)$, and b_n and b'_n are not equal unless $\lambda = \frac{1}{2}$.

²By Bayes' rule, as $b_n \rightarrow \lambda$, $\Pr(\theta = H|X_n = 1) = \lambda$ and $\Pr(\theta = L|X_n = 0) = 1 - \lambda$, and so the prediction X_n is not informative at all; at $b_n = 1$, $\Pr(\theta = H|X_n = 1) = \Pr(\theta = L|X_n = 0) = 1$, and so X_n fully reveals θ .

³It can be shown that as b_n increases, X_n becomes more informative in the Blackwell sense (Blackwell, 1953). Throughout the paper, when we say "more accurate," we mean more Blackwell informative.

the machine and human are both mistaken in their predictions is negligible.

Assumption 2. Conditional on θ , the probability of both X_m and X_h being mistaken is zero, i.e., $\Pr(X_m = 0, X_h = 0|\theta = H) = \Pr(X_m = 1, X_h = 1|\theta = L) = 0$.

Given our context and research questions, Assumption 2 is not overly restrictive because it implies $\Pr(\theta = H|X_m = 1, X_h = 1) = \Pr(\theta = L|X_m = 0, X_h = 0) = 1$, that is, the market state is fully revealed when the machine and human forecasts agree with each other. This reflects the idea that the best forecasts often result from combining human and machine insights (Ibrahim et al., 2021) and is somewhat similar to the information structure assumed in Boyacı et al. (2024).⁴

Assumptions 1 and 2, along with the machine and human accuracies b_m and b_h defined earlier, result in the following conditional probability distribution of (X_m, X_h) :

$$\left\{ \begin{array}{l} p_{11|H} = b_m + b_h - 1, \\ p_{10|H} = 1 - b_h, \\ p_{01|H} = 1 - b_m, \\ p_{00|H} = 0 \end{array} \right. \quad \text{and} \quad \left\{ \begin{array}{l} p_{11|L} = 0, \\ p_{10|L} = \frac{\lambda}{1-\lambda}(1 - b_m), \\ p_{01|L} = \frac{\lambda}{1-\lambda}(1 - b_h), \\ p_{00|L} = 1 - \frac{\lambda}{1-\lambda}(1 - b_m) - \frac{\lambda}{1-\lambda}(1 - b_h). \end{array} \right. \quad (1)$$

Using the above conditional probabilities and Bayes' rule, we can derive the unconditional probabilities of each realization of (X_m, X_h) , denoted by $p_{ij} = \Pr(X_m = i, X_h = j)$ for $i, j \in \{0, 1\}$, and the resulting posterior probabilities of $\theta = H$ conditional on (X_m, X_h) , as summarized in Table 1. As an aside, we note that the relative accuracy metric β defined above is also equal to the probability of the machine forecast X_m being correct whenever the predictions X_m and X_h diverge.

Table 1: Unconditional probability distribution of (X_m, X_h) and the posterior probabilities of $\theta = H$.

X_m	X_h	$\Pr(X_m, X_h)$	$\Pr(\theta = H X_m, X_h)$
1	1	$p_{11} = \lambda(b_m + b_h - 1)$	1
1	0	$p_{10} = \lambda(2 - b_m - b_h)$	β
0	1	$p_{01} = \lambda(2 - b_m - b_h)$	$1 - \beta$
0	0	$p_{00} = 1 - \lambda(3 - b_m - b_h)$	0

Finally, given the possibility of overlapping information feeding the machine and human predictions, it would seem unrealistic to allow the unconditional predictions X_m and X_h to be negatively

⁴Motivated primarily by medical decisions, Boyacı et al. (2024) consider a setting in which the true state is good iff both the machine's and human's evaluations are good. In this regard, our Assumption 2 is in line with their model. However, in their model, a negative machine evaluation will fully reveal that the true state is bad. This fits well with healthcare contexts where decisions are critical. Differently, we consider demand forecasting, and therefore, in our model the machine forecast X_m alone will not fully the true state, regardless of whether X_m is positive or negative.

correlated. Therefore, in all that follows, we restrict attention to independent or positively correlated predictions. X_m and X_h have a correlation coefficient $\rho(X_m, X_h) \geq 0$ iff $b_m + b_h \geq 1 + \lambda$.⁵ Therefore, in this paper, we only consider the parameter space $\{(b_m, b_h) \in (\lambda, 1]^2 : b_m + b_h \geq 1 + \lambda\}$. All the probabilities in (1) are well-defined within this space.

Human Prediction Effort. The manager must exert costly effort to generate a market-state prediction. In the base model, we assume that the manager chooses between two effort levels, $e \in \{0, 1\}$. If $e = 0$, they do not obtain any information about X_h . If $e = 1$, they acquire a signal s_h that fully reveals the prediction X_h , i.e., $s_h = X_h$, while incurring a non-monetary effort cost $k > 0$. The effort cost κ captures the time and attention required by the manager to gather and process information to form the prediction X_h . A binary effort choice is often adopted in the literature to model information acquisition (Taylor and Xiao, 2009; Tang and Girotra, 2017). However, in §5.4, we allow the manager to choose a continuous effort level $e \in [0, 1]$ where s_h partially reveals X_h if $e < 1$. We show that our main findings based on binary effort continue to hold. Although $s_h = X_h$ in the base model, we choose to work with the notation s_h (instead of replacing s_h everywhere with X_h) to more easily allow for the later generalization.

Machine Information Provision. The machine communicates the value of its prediction X_m to the manager through a message (signal) realization $s_m \in \Gamma$, where Γ denotes the set of message realizations. We consider a binary message space $\Gamma = \{0, 1\}$ while noting that restricting to a binary space is without loss.⁶ We use η to denote the messaging policy of the machine, which is formally a mapping from any realized value of the machine prediction X_m to a distribution over message realizations. That is, $\eta : \{0, 1\} \rightarrow \Delta(\Gamma)$, where $\Delta(\Gamma)$ denotes the set of all probability distributions on support $\Gamma = \{0, 1\}$. Specifically, let $\eta(s_m|X_m)$ denote the probability of the machine message realization being s_m conditional on a realized prediction X_m . For instance, a machine that fully reveals the value of X_m has $\eta(1|1) = 1$ and $\eta(0|0) = 1$. In general, however, $\eta(1|1)$ and $\eta(0|0)$ need not be one, in which case the messaging policy only partially reveals the machine prediction X_m . As described below, the firm determines the messaging policy.

The Manager's Decisions. Equipped with a machine, the manager solves a two-stage problem: prediction effort level followed by a quantity decision. First, upon receiving a machine message s_m , they decide to either exert no effort and rely exclusively on the machine prediction or to exert

⁵The correlation coefficient is given by $\rho(X_m, X_h) = \frac{p_{11} - (p_{01} + p_{11})(p_{10} + p_{11})}{\sqrt{(p_{00} + p_{01})(p_{00} + p_{10})(p_{01} + p_{11})(p_{10} + p_{11})}}$ where the unconditional probabilities are given in Table 1.

⁶When solving for the optimal messaging policy, we do not require any restriction on the message space and show that binary messages are sufficient to achieve optimality.

effort at a cost k to obtain their own additional signal $s_h = X_h$. Next, conditional on the available information, a single signal s_m in case of $e = 0$ or a pair of signals s_m and s_h in case of $e = 1$, the manager chooses an order quantity Q to maximize the expected profit based on their posterior belief about the market state θ . This second decision is a newsvendor problem.

Let $\mathbf{s}$ be the information set available to the manager. Let us write $\mathbf{s} = (s_m, \emptyset)$ if they opt not to elicit s_h , and $\mathbf{s} = (s_m, s_h)$ otherwise. Given any $\mathbf{s}$, the manager chooses an order quantity Q to maximize the following profit function:

$$\pi(Q|\mathbf{s}) = r\mathbb{E}_\theta[\min(a_\theta, Q)|\mathbf{s}] - cQ \quad (2)$$

where the expectation is taken over market state θ based on the manager's posterior belief formed from information set $\mathbf{s}$.

Define $\Pi(\mathbf{s}) = \max_{Q \geq 0} \pi(Q|\mathbf{s})$ as the optimal expected profit conditional on information set $\mathbf{s}$. Based on the realized machine signal s_m , the manager decides on their prediction effort e . If no effort is exerted ($e = 0$), then the manager will make the ordering decision based on information set $\mathbf{s} = (s_m, \emptyset)$, resulting in an expected profit $\Pi(s_m, \emptyset)$. If choosing to exert effort ($e = 1$), the manager will have access to both signals (s_m, s_h) , leading to an expected profit $\Pi(s_m, s_h)$. Because s_h is random when deciding the effort e , the manager's expected payoff for $e = 1$ is given by $\sum_{s_h \in \{0,1\}} \Pr(s_h|s_m) \Pi(s_m, s_h) - k$, where $\Pr(s_h|s_m)$ represents the probability of the human-generated signal being s_h conditional on machine signal realization s_m . The manager chooses the effort level that leads to the maximum expected payoff. Therefore, the optimal effort level is given as follows.

$$e^*(s_m) = \begin{cases} 1 & \text{if } \sum_{s_h \in \{0,1\}} \Pr(s_h|s_m) \Pi(s_m, s_h) - k \geq \Pi(s_m, \emptyset) \\ 0 & \text{otherwise.} \end{cases} \quad (3)$$

In formulating the manager's effort and quantity decision objectives, we have scaled the non-monetary prediction-effort cost k in monetary units and have assumed that the manager's compensation is structured so that they choose an order quantity to maximize a scalar multiple of the expected monetary profit (with the scalar set to one for notational ease). As discussed and expanded on in Remark 1 below, this is equivalent to assuming the manager is compensated through a performance-based linear contract. The contract is assumed exogenous in our base model but we explore contract-parameter optimization in §5.1.

The Firm's Decision on Machine Information Provision. The firm, having decided

to provide a prediction machine to the manager, determines the machine’s messaging policy so as to maximize its expected profit. The firm cannot internalize the manager’s non-monetary effort cost k because the time and cognitive effort that a human spends processing information is usually unverifiable or non-contractible. Therefore, the manager’s prediction effort cost induces a misalignment in incentives between the firm and the manager.⁷ We intentionally focus on an informational approach to managing this misalignment and we therefore exclude other approaches such as outcome-contingent compensation in which the manager receives monetary payments contingent on the realized machine message and profit.

The firm configures (“tunes”) the prediction machine by determining a messaging policy η to maximize the expected *monetary* profit. The firm’s problem can be formulated as follows.

$$\max_{\eta} \mathbb{E}_{X_m \sim \lambda} \mathbb{E}_{s_m \sim \eta(\cdot|X_m)} [e^*(s_m) \mathbb{E}_{s_h|s_m} \Pi(s_m, s_h) + (1 - e^*(s_m)) \Pi(s_m, \emptyset)] \quad (4)$$

where $e^*(s_m)$ is given by (3). Recall that η , as the decision variable of (4), is a collection of the probability distributions $\eta(s_m|X_m)$ for all $X_m \in \{0, 1\}$.

Remark 1. Our model has assumed, in effect, that the manager is compensated through a linear performance-based contract. That is, the manager receives $w_B + w_P \pi$, where w_B is the base salary, w_P is the commission rate, and π is the firm’s realized monetary profit not including compensation. Linear (commission-based) compensation contracts are commonly used in practice and widely adopted in the literature exploring managerial compensation (see, e.g., Jin, 2002; Baiman et al., 2010; Long and Nasiry, 2020). Let the unscaled managerial effort cost be denoted by $\kappa > 0$. Given w_B and w_P , the manager maximizes (with the relevant expectation) $w_B + w_P \pi - \kappa e$, which is equivalent to optimizing $\pi - ke$ by defining a scaled effort cost $k = \frac{\kappa}{w_P}$. In the base model, we focus on optimizing the machine information provision while assuming the compensation parameters $\{w_B, w_P\}$ to be fixed.⁸ Thus, without loss, we work with the scaled effort cost k and the transformed manager payoff $\pi - ke$ hereafter until compensation design is explored in §5.1.

Remark 2. Although we formulate the problem as choosing among messaging policies η that

⁷Our main results can be extended if part of the human’s information acquisition cost can be internalized by the organization. For example, if the human decision maker, after observing the machine’s message, decides to pay a third party for a market survey, then we can readily incorporate the monetary cost for the survey in the machine designer’s objective; however, the additional time and attention required to process the survey results are still incurred by the human.

⁸Changing existing compensation terms because of AI adoption may be costly in practice because it may be subject to mandatory bargaining (Young and Morris, 2023) and/or may result in a long-term cost of losing productive employees (Sandvik et al., 2021).

partially or fully reveal X_m , an alternative interpretation of our problem is that the firm can choose (“tune”) the machine’s true positive and true negative prediction probabilities subject to a set of constraints governed by b_m . To see this, note that messaging policies $\eta(s_m|X_m)$ conditional on X_m can be readily transformed to probability distributions conditional on the market state θ , denoted by $\tilde{\eta}(s_m|\theta)$. That is, $\tilde{\eta}(s_m|\theta) = \sum_{X_m \in \{0,1\}} \eta(s_m|X_m) \Pr(X_m|\theta)$, with $\tilde{\eta}(1|H)$ and $\tilde{\eta}(0|L)$ being the true positive and true negative probabilities, respectively. When optimizing over η , we are essentially choosing among machine information structures $\tilde{\eta}$ that are Blackwell (weakly) less informative than the information structure obtained when X_m is fully disclosed (Blackwell, 1953). With that interpretation, the tuned machine fully discloses its signal to the manager with information structure $\tilde{\eta}$. As will be established later, partial disclosure of X_m can be optimal for problem (4), which indicates that the firm may be better off tuning the machine to be less informative than its maximum prediction capability (defined by b_m) even when higher informativeness (subject to an upper limit) is costless.

Remark 3. In formulating Problem (4), we have assumed that the firm can credibly announce and commit to a messaging policy η and that the manager is aware of η and makes the prediction-effort and ordering decisions based on the realized machine message s_m . The former commitment assumption is reasonable in our context because a machine’s precision can be quantified and communicated beforehand. The latter assumption, under which the manager exerts prediction effort after receiving the machine message, is consistent with the human-machine literature (e.g., Boyacı et al., 2024). It is also justified by the temporal flexibility of managerial effort: even if the manager is required to submit an initial prediction before observing the machine signal, they can still withhold part of their effort until after seeing s_m . As discussed earlier, our model reflects a setting in which the manager retains responsibility for the order-quantity decision, rather than being forced to delegate it to the machine. Hence, even if the manager must first submit a preliminary prediction, they can still refine it after observing the machine signal.

4 Results

4.1 Optimal Order Quantity

Given any realized information set $\mathbf{s}$, regardless of whether the manager has generated their own signal s_h or not, the optimal order quantity is influenced by $\mathbf{s}$ only through their posterior belief about the market state, which is formed based on $\mathbf{s}$. Let μ denote the posterior belief that the

market state $\theta = H$, that is, $\mu = \Pr(\theta = H|\mathbf{s})$. Note that μ is affected by the machine's accuracy b_m and signal structure $\boldsymbol{\eta}$, by the realized machine signal s_m , and by the realized human signal s_h and accuracy b_h (if effort is exerted).

To begin the analysis, we first derive the optimal order quantity for any given posterior belief μ . Define $\phi = c/r$ as the cost-to-revenue ratio of the product. Note that $1 - \phi$ is commonly referred to as the critical fractile in newsvendor problems. A lower ϕ represents a product with a higher profit margin.

Lemma 1. *Given any posterior belief $\mu = \Pr(\theta = H|\mathbf{s})$, the optimal order quantity is given by $Q^*(\mu) = a_L$ if $\mu \leq \phi$ and $Q^*(\mu) = a_H$ if $\mu > \phi$. The resulting expected profit can be expressed as a function of μ :*

$$\pi^*(\mu) = \begin{cases} (r - c)a_L & \text{if } \mu \leq \phi, \\ (r - c)a_L + (r\mu - c)(a_H - a_L) & \text{if } \mu > \phi. \end{cases} \quad (5)$$

In our two-point demand setting, it is clearly optimal to order either a_L units or a_H units. Ordering a_L guarantees there is no leftover inventory in either state but forgoes some demand if the market state is high. Ordering a_H guarantees demand is always met but risks leftover inventory (overstocked) if the market state is low. Intuitively, the more the manager believes the market state to be high, the more willing they should be to bear the overstock risk. Lemma 1 establishes that there exists a threshold such that ordering a_H units is optimal iff the posterior belief μ is above the threshold, and that the threshold is exactly the cost-to-revenue ratio ϕ .

4.2 Benchmarks: No Machine and Machine with Full Information Disclosure

We begin by analyzing two benchmark settings: one in which the firm does not provide a prediction machine to the manager, and one in which a machine is provided and the machine fully discloses its prediction to the manager. In the no-machine setting, the manager chooses whether to acquire additional information s_h based only on the prior λ . In the full-disclosure machine setting, the manager chooses their prediction effort based on the realized machine message $s_m \in \{0, 1\}$ that fully reveals $X_m \in \{0, 1\}$, which is equivalent to having a messaging policy $\eta(1|1) = \eta(0|0) = 1$ and $\eta(0|1) = \eta(1|0) = 0$.

To characterize the optimal prediction effort in both settings, we define $l_n = \frac{\hat{k} + \lambda(1 - b_h)}{1 - \lambda}$, $u_n = b_h - \frac{\hat{k}}{\lambda}$, $l_0 = \frac{(1 - \lambda)\hat{k}}{1 - \lambda(3 - b_m - b_h)}$, $u_0 = \frac{\lambda(1 - b_m) - (1 - \lambda)\hat{k}}{\lambda(2 - b_m - b_h)}$, $l_1 = \frac{\hat{k} + 1 - b_h}{2 - b_h - b_m}$, and $u_1 = 1 - \frac{\hat{k}}{b_m + b_h - 1}$, where $\hat{k} = \frac{k}{r(a_H - a_L)}$ scales the effort cost relative to the potential monetary gain in the high state. Proposition 1 below establishes the conditions under which the manager will exert effort with and without a machine.

These conditions can be represented by three triangular regions in the k - ϕ coordinate plane, as illustrated in Figure 1.

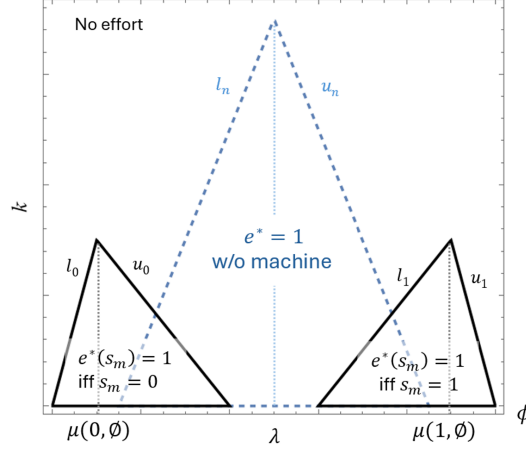


Figure 1: Optimal effort e^* without machine and with full-information machine

Proposition 1. (i) Without machine information, the manager exerts effort $e^* = 1$ iff $\phi \in [l_n, u_n]$.
(ii) With full machine information, contingent on the realization of machine message $s_m \in \{0, 1\}$, the manager's optimal effort level $e^*(s_m)$ is given by $e^*(0) = 1$ iff $\phi \in [l_0, u_0]$, and $e^*(1) = 1$ iff $\phi \in [l_1, u_1]$.

Without a machine, the manager exerts effort iff ϕ falls into an interval $[l_n, u_n]$, represented by the dashed triangular region which centers around the prior λ . Recall from Lemma 1 that the optimal order quantity depends crucially on whether the belief that demand is high ($\theta = H$) is greater or less than ϕ . If no effort is exerted, then this belief is the prior λ , and the manager orders a_H (a_L) if $\phi \leq \lambda$ ($\phi > \lambda$). Exerting effort (to generate a prediction) will shift the posterior belief μ to the left or right of the prior. From Lemma 1, a shift to the left (right) changes the order quantity from a_H to a_L (from a_L to a_H) only if the posterior $\mu \leq \phi$ ($\mu > \phi$). In the relevant range of ϕ to the left (right) of λ , the benefit of exerting effort is the avoidance of leftover inventory (the gain from higher sales), and this benefit is increasing (decreasing) in ϕ for a fixed r . It is optimal to exert effort only when this benefit outweighs the effort cost k , and so we have a triangular region centered around the prior.

When a full-disclosure machine is provided, the manager chooses their prediction effort based on their updated belief according to the realized machine message s_m . Denote by $\mu(s_m, \emptyset)$ the updated belief that $\theta = H$ upon observing s_m (but without acquiring s_h); we have $\mu(0, \emptyset) < \lambda < \mu(1, \emptyset)$. Two effort-exerting intervals, $[l_0, u_0]$ and $[l_1, u_1]$ emerge for $s_m = 0$ and $s_m = 1$, respectively; see

the solid triangular regions in Figure 1. The effort-exerting region for $s_m = 0$ is on the lower end of the ϕ axis, whereas the region for $s_m = 1$ is on the higher end. We explain each in turn.

A low ϕ signifies a high profit margin, and so the manager is somewhat inclined to order for the high market state, i.e., order a_H and take the overstocking risk instead of the unmet demand risk. A high-state prediction from the machine ($s_m = 1$) aligns with this inclination, and so further refining the posterior through prediction effort does not generate much value. In contrast, a low-state prediction from the machine ($s_m = 0$), pushes back on the manager’s “order high” inclination, and so generating a prediction to further refine the posterior helps the manager decide whether it should go with its inclination (order high) or adjust based on the machine signal (which directionally pushes towards ordering low). A high ϕ signifies a low profit margin, and so the manager is somewhat inclined to order for the low market state, i.e., order a_L and take the unmet demand risk instead of the overstocking risk. Analogous to above, it is in the manager’s interest to exert prediction effort only when the machine signal runs counter to this inclination to order low; that is, it is optimal to exert effort only if the machine signal is high ($s_m = 1$). If one thinks of the new prior (post machine prediction) as being the prior for the effort exertion decision, then the logic described above to explain the dashed triangular region for the no-machine case applies also to the with-machine case, but the relevant prior is now to the left of λ if $s_m = 0$ or to the right if $s_m = 1$.

Crucially, Proposition 1 establishes that providing a prediction machine can dampen the manager’s prediction effort. For example, for those (ϕ, k) pairs in Figure 1 lying within the large dashed triangle but not within either solid triangle, the manager never exerts effort if given a (full-disclosure) machine but would exert effort in the absence of a machine. An important question, therefore, and the one we now explore, is whether the firm can mitigate this effort-dampening impact by tuning the machine’s messaging policy.

4.3 Optimal Information Disclosure

We now consider the firm’s problem given by (4) and determine the optimal messaging policy η^* that maximizes the expected monetary profit. For any given messaging policy η , the machine message s_m influences the optimal effort $e^*(s_m)$ solely through the resulting posterior belief that the machine prediction $X_m = 1$. Per Bayes’ rule, $\Pr(X_m = 1|s_m) = \frac{\eta(s_m|1)\lambda}{\eta(s_m|1)\lambda + \eta(s_m|0)(1-\lambda)}$. Let v denote the posterior probability that $X_m = 1$. Note that v is a random variable ex ante because X_m and s_m are uncertain. The information design literature has shown that optimizing over the messaging policy is equivalent to determining a particular Bayes-plausible distribution of posteriors (Kamenica

and Gentzkow, 2011). Following this approach, we will hereafter work with the posterior v , instead of $\eta(s_m|X_m)$. An important feature of our problem is that the firm controls the information regarding the machine prediction X_m , which itself is only an estimator of the market state θ . Our analysis will leverage the posterior of X_m , rather than the posterior of θ .

4.3.1 Optimal Prediction Effort Given a full-disclosure machine, that is, $\eta(1|1) = \eta(0|0) = 1$, the posterior v will be either one (with probability λ) or zero (with probability $1 - \lambda$). Nevertheless, if one allows partial disclosure such that $\eta(1|1) < 1$ or $\eta(0|0) < 1$, then v could lie strictly between zero and one. This motivates us to characterize the manager's optimal effort generally as a function of $v \in [0, 1]$. To do so, we establish a mapping from v , the posterior belief regarding X_m , to μ , the posterior regarding θ . Then, for any posterior $v \in [0, 1]$, we rewrite the expected monetary profits with and without generating the human prediction as $\Pi_I(v)$ and $\Pi_N(v)$, respectively. Consequently, the optimal prediction effort can be recast as $e^*(v) = 1$ iff $\Pi_I(v) - \Pi_N(v) \geq k$. Technical details about this transformation are relegated to Appendix A.

The optimal prediction effort $e^*(v)$ can be characterized using the two thresholds defined below.

$$\psi_1 = \frac{\phi\lambda(2 - b_m - b_h) - \lambda(1 - b_m) + (1 - \lambda)\hat{k}}{(1 - \phi)(b_m + b_h - 1 - \lambda) + \lambda(1 - b_h)} \quad (6)$$

and

$$\psi_2 = \frac{\phi(1 - \lambda(3 - b_m - b_h)) - (1 - \lambda)\hat{k}}{\phi(b_m + b_h - 1 - \lambda) + (1 - \lambda)(1 - b_h)}. \quad (7)$$

Lemma 2. *For any given posterior belief $v = \Pr(X_m = 1|s_m)$, the manager's optimal prediction effort $e^*(v)$ can be characterized as follows. If $\hat{k} > \bar{K}$, then $e^*(v) = 0$ for all v (Case (N)). Otherwise, there exists two thresholds $\psi_1 \leq \psi_2$ such that $e^*(v)$ is given by one of the following three cases: (i) for $\phi \leq u_0$, $e^*(v) = 1$ iff $v \leq \psi_2$; (ii) for $\phi \geq l_1$, $e^*(v) = 1$ iff $v \geq \psi_1$; (iii) for $u_0 < \phi < l_1$, $e^*(v) = 1$ iff $\psi_1 \leq v \leq \psi_2$. Moreover, for each $i = 1, 2$, provided $0 < \psi_i < 1$, ψ_i is increasing in ϕ .*

Lemma 2 shows that the optimal prediction effort $e^*(v)$ depends on $\hat{k} = k/(r(a_H - a_L))$ and on the cost-to-revenue ratio $\phi = c/r$. The structure of $e^*(v)$ is illustrated in Figure 2. If $\hat{k}$ is higher than a threshold $\bar{K}$, the manager will never exert effort to make a prediction, regardless of the posterior v ; this is Case (N) of Lemma 2.⁹ Note that the threshold depends on ϕ . When ϕ is close to 0 or 1, the value of demand information is diminished because, regardless of the human prediction, the optimal order quantity is a_H (a_L) because of the extremely low (high) procurement

⁹The closed-form expression of $\bar{K}$ can be found in the proof of Lemma 2.

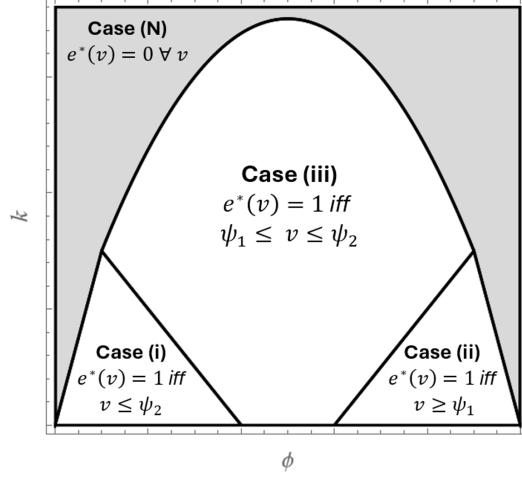


Figure 2: Structure of the optimal effort $e^*(v)$ for different combinations of k and ϕ (Parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$)

cost relative to the selling price. As a result, the parameter region for Case (N) is enlarged as ϕ becomes extreme, as shown in Figure 2. If $\hat{k}$ is below $\bar{K}$, the optimal effort is posterior-dependent. It can be shown that $\Pi_I(v) - \Pi_N(v)$ is unimodal in v . As a result, $e^*(v) = 1$ iff v falls in an interval $[\psi_1, \psi_2]$, where $\psi_1 \leq 0$ for sufficiently small ϕ and $\psi_2 \geq 1$ for sufficiently large ϕ . Thus, Cases (i) - (iii) of Lemma 2 follow.

Cases (i) and (ii) correspond exactly to the two solid triangular regions presented earlier in Figure 1 where, given a full-disclosure machine, the manager exerts effort only when receiving a low message ($s_m = 0$) and a high message ($s_m = 1$), respectively. Recall that the low and high messages from a full-disclosure machine must imply $v = 0$ and $v = 1$. Now, Lemma 2 presents a more general result: For high-margin products (small ϕ), the manager exerts effort if the posterior v is lower than ψ_2 ; for low-margin products (large ϕ), the manager exerts effort if the posterior v is higher than ψ_1 .

More importantly, compared to the full-disclosure benchmark, a new case – Case (iii) – emerges when we consider any posterior $v \in [0, 1]$. With a full-disclosure machine, the manager never exerts effort for moderate profit margins or $u_0 < \phi < l_1$. Nevertheless, Case (iii) of Lemma 2 indicates that the manager is incentivized to exert effort if their posterior v lies in an interval $[\psi_1, \psi_2]$. Moreover, we can establish that both ψ_1 and ψ_2 increase with ϕ , that is, the interval $[\psi_1, \psi_2]$ shifts toward the higher end of v as the profit margin becomes lower, reflecting the same intuition as before that the manager has greater incentive to exert prediction effort when the machine hints at high demand for low-margin products or when the machine hints at low demand for high-margin products.

4.3.2 Optimal Messaging Policies Lemma 2 implies that the firm might benefit from motivating effort by obfuscating the machine prediction X_m so that the manager's posterior belief v falls into the interval $[\psi_1, \psi_2]$. However, obfuscation also compromises to some extent the value of the machine's information. This tradeoff leads to the following important question: when it is optimal to obfuscate X_m and to what degree it should be obfuscated. When $\hat{k} > \bar{K}$, effort cannot be induced for any posterior v (Lemma 2 Case (N)) and so full disclosure is trivially optimal. Thus, without loss, we will assume $\hat{k} \leq \bar{K}$ hereafter.

Lemma 2 implies that the firm might benefit from motivating effort by obfuscating the machine prediction X_m such that v falls into the interval $[\psi_1, \psi_2]$. Important questions are when it is optimal to do so and to what extent X_m should be obfuscated, because obfuscation will, on the other hand, sacrifice some value of the machine's information. Clearly, when $k > \bar{K}$ such that effort cannot be induced for any posterior, full disclosure is trivially optimal. Thus, without loss, we will assume $\hat{k} \leq \bar{K}$ hereafter.

Following the information design literature (Kamenica and Gentzkow, 2011), the optimal expected payoff is achieved on the upper concave envelope of the expected monetary profit $\Pi(v)$, where $\Pi(v) = \Pi_I(v)$ if $e^*(v) = 1$, and $\Pi(v) = \Pi_N(v)$ otherwise. We will leverage this property to derive the optimal messaging policy η^* that solves (4). Throughout the paper, in cases where multiple information structures are optimal for problem (4), we break the tie by choosing the optimal messaging policy η^* that is more informative.

If the machine prediction is fully disclosed, it then follows from our earlier results that the manager only exerts effort in Case (i) if given a low message $s_m = 0$, only exerts effort in Case (ii) if given a high message $s_m = 1$, and exerts no effort in Case (iii). In Proposition 2, we characterize the optimal messaging policy η^* under each of these three cases. To simplify notation, we will write $\eta^*(1|1)$ and $\eta^*(0|0)$ as η_1^* and η_0^* , representing the probability of the machine message accurately communicating X_m ; accordingly, $\eta^*(0|1) = 1 - \eta_1^*$ and $\eta^*(1|0) = 1 - \eta_0^*$.

Proposition 2. *The optimal messaging policies are characterized as follows.*

(i) If $\phi \leq u_0$, full information disclosure **(F)** is optimal, that is,

$$\eta_1^* = \eta_0^* = 1 \tag{F}$$

(ii) If $\phi \geq l_1$, full information disclosure **(F)** is optimal.

(iii) If $u_0 < \phi < l_1$, there exist thresholds $\hat{\phi}_1$ and $\hat{\phi}_2$ such that (a) for $u_0 < \phi < \hat{\phi}_1$, partial information disclosure is optimal with messaging policy **(D)** given below

$$\eta_1^* = \frac{\lambda - \psi_1}{1 - \psi_1} \frac{1}{\lambda} \text{ and } \eta_0^* = 1. \quad (\text{D})$$

(b) for $\hat{\phi}_1 \leq \phi \leq \hat{\phi}_2$, full information disclosure **(F)** is optimal.

(c) for $\hat{\phi}_2 < \phi < l_1$, partial information disclosure is optimal with messaging policy **(E)** given below

$$\eta_1^* = 1 \text{ and } \eta_0^* = \left(1 - \frac{\lambda}{\psi_2}\right) \frac{1}{1 - \lambda}. \quad (\text{E})$$

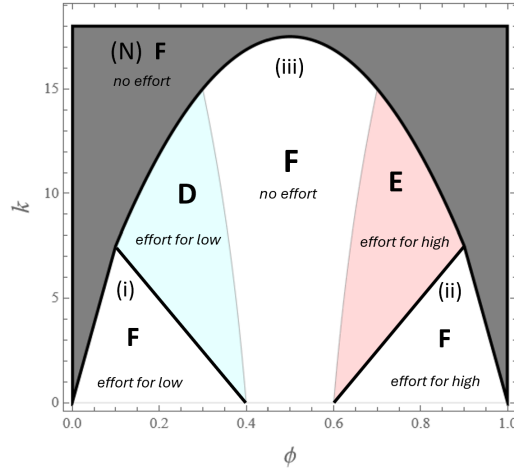


Figure 3: Structure of the optimal messaging policies as in Proposition 2 (Parameter values: $b_m = 0.9$, $b_h = 0.85$, $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$, and c is varied based on ϕ .)

Proposition 2 is numerically illustrated in Figure 3. First, note that $\hat{k} > \bar{K}$ in the dark region (N) and therefore full disclosure is trivially optimal. Proposition 2 focuses on the optimal messaging policies for the rest of the parameter space.

In Cases (i) and (ii), full disclosure is optimal. That is, if the manager already has the incentive to exert effort for one of the signal realizations under full machine information (for the low signal in case (i) and the high signal in case (ii)), then partial disclosure does not add value and the firm should deploy a full-disclosure machine. Put differently, the full-disclosure triangles in Figure 3 are the same as the signal-dependent effort triangles in Figure 1.

In Case (iii), where the manager never exerts effort under full machine information, partially disclosing the machine prediction can be optimal, as we now discuss.

For $\phi < \hat{\phi}_1$, the optimal messaging policy is given by (D), where (D) stands for “downplay.” The machine always communicates a low prediction by sending a low message (i.e., $\eta_0^* = 1$), but it probabilistically downplays a high prediction by sending a low message with positive probability $\eta^*(0|1) = 1 - \eta_1^* > 0$. As established earlier, for high-margin products, the manager is inclined to exert effort when receiving a low prediction. In region (D) of Figure 3, the profit margin is high but not high enough to induce prediction effort under a full-disclosure machine ($u_0 < \phi < \hat{\phi}_1$). Therefore, downplaying a high prediction to a low one (with some probability) is optimal because it motivates managerial effort under a low machine message. As shown graphically in Figure 3, doing so expands the effort-exerting region for low messages.

For $\phi > \hat{\phi}_2$, the optimal messaging policy is given by (E), where (E) stands for “exaggerate.” The machine always communicates a high prediction by sending a high message (i.e., $\eta_1^* = 1$), but it probabilistically exaggerates a low prediction by sending a high message with positive probability $\eta^*(1|0) = 1 - \eta_0^* > 0$. For low-margin products, the manager is inclined to exert effort given a high prediction. In region (E) of Figure 3, the profit margin is low but not low enough to induce prediction effort under a full-disclosure machine ($\hat{\phi}_2 < \phi < l_1$). Therefore, exaggerating a low prediction to a high one (with some probability) can be optimal because it motivates managerial effort under a high machine messages; thus, expanding the effort-exerting region for high messages in Figure 3.

For ϕ within the medium range $[\hat{\phi}_1, \hat{\phi}_2]$, inducing managerial effort requires so much obfuscation (either downplaying or exaggerating) that the loss of machine information cannot be justified by the benefit resulting from the manager’s prediction. Consequently, the optimal policy reverts to full disclosure without inducing human effort.

Next, to illustrate how partial disclosure may lead to a higher profit than full disclosure, we consider a numerical example in which $b_m = 0.9$, $b_h = 0.85$, $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$ and $k = 8$. Given $c = 8$ and hence $\phi = 0.8$, by Lemma 2, the manager will exert effort iff the posterior v falls between $\psi_1 = 0.72$ and $\psi_2 = 0.9455$. As depicted by the black solid line in Figure 4, the expected monetary profit $\Pi(v)$ as a function of the posterior v consists of three pieces: for $v \in [0, \psi_1)$, $\Pi(v)$ is constant because the manager does not acquire s_h and orders the low quantity of a_L units; for $v \in [\psi_1, \psi_2]$, the manager acquires s_h and adjusts their order size based on the realization of s_h , and so the monetary profit $\Pi(v)$ jumps upward and increases in v ; for $v \in (\psi_2, 1]$, the manager does not acquire s_h but orders a_H units, and so $\Pi(v)$ increases in v .

Under full disclosure, the machine message s_m precisely reveals X_m . Consequently, $v_F = 1$ and

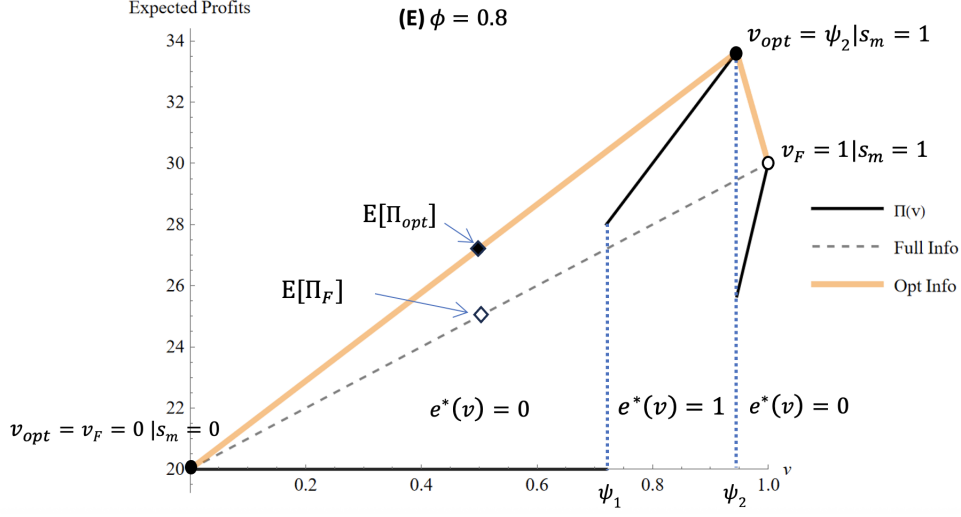


Figure 4: Illustration of the expected monetary profits under the optimal messaging policy (E) in Proposition 2(iii) [Parameter values: $\beta = 0.6$ with $b_m = 0.9$ and $b_h = 0.85$, $k = 8$, $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$ and $c = 8$]

$\Pi(v_F) = 30$ if $s_m = 1$ (which occurs with probability 0.5), whereas $v_F = 0$ and $\Pi(v_F) = 20$ if $s_m = 0$ (probability 0.5). The manager exerts no effort in either scenario, resulting an expected profit $\Pi_F = 25$, as marked by the middle point of the dashed line in Figure 4.

Given $\phi = 0.8$, policy (E) is optimal with $\eta_1^* = 1$ and $\eta_0^* = 0.942$, whereby the machine always communicates high predictions as high, but may exaggerate low predictions as high with probability $1 - \eta_0^* = 0.058$. Consequently, the manager forms the posterior belief as follows. If receiving a low message $s_m = 0$, they will ascertain a low prediction $X_m = 0$, i.e., $v_{opt} = 0$; however, if receiving a high message $s_m = 1$, they are aware that the actual prediction may possibly be low with probability 0.058, and applying Bayes' rule leads to a posterior $v_{opt} = 0.9455$, which coincides with ψ_2 . Note that $v_{opt} = \psi_2$ provides just enough incentive for the manager to exert effort. As such, the optimal policy induces effort (under a high message) with *minimal* obfuscation of the machine prediction. This partial disclosure policy results in an expected profit $\Pi_{opt} = 27.21$, an 8.8% improvement over the expected profit under full disclosure $\Pi_F = 25$.

Figure 5 illustrates the two other policies in Proposition 2(iii) that arise for different ϕ . Different to $\phi = 0.8$, the product has a high profit margin when $\phi = 0.2$, and from Lemma 2(iii), the manager exerts prediction effort only when v falls into a lower interval given by $\psi_1 = 0.0545$ and $\psi_2 = 0.28$. In this case, policy (D) is optimal under which the machine precisely reveals low predictions while downplaying high predictions as low with probability $1 - \eta_1^* = 0.058$. Consequently, a low message retains just enough probability that the actual prediction is high to induce effort, that is, $v_{opt} = \psi_1$.

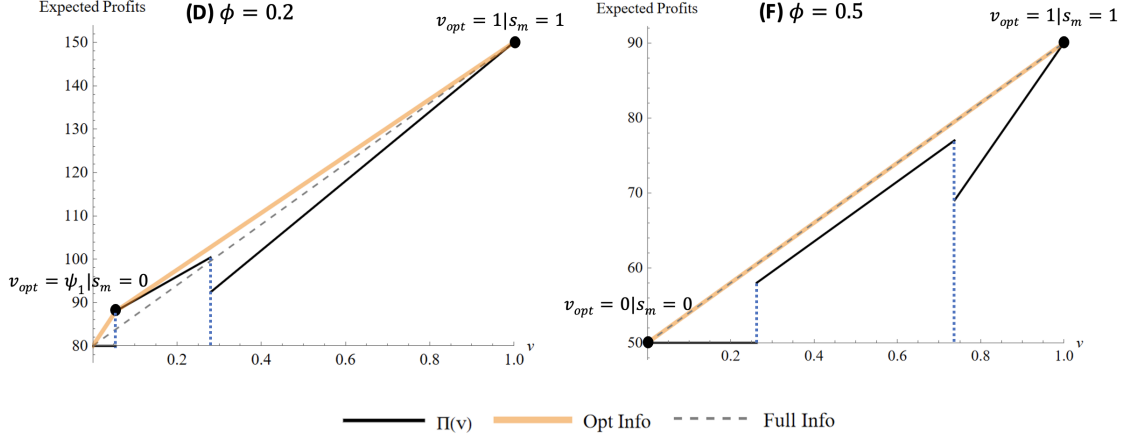


Figure 5: Illustration of the expected monetary profits under the optimal messaging policies (D) and (F) in Proposition 2(iii) [Parameter values: $\beta = 0.6$ with $b_m = 0.9$ and $b_h = 0.85$, $k = 8$, $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$ and c varied according to ϕ given in the graph]

See the left panel of Figure 5. When $\phi = 0.5$, the manager exerts effort only when v falls into a moderate interval given by: $\psi_1 = 0.2625$ and $\psi_2 = 0.7375$. Inducing a posterior $v = \psi_1$ or ψ_2 comes at the cost of withholding too much machine information, and so fully disclosing the machine prediction and not motivating effort becomes optimal. See the right panel of Figure 5. In Appendix C, we present a more comprehensive numerical study showing that partially disclosing machine information can generate significantly larger profit gains than a full-disclosure machine.

Our analysis highlights a tradeoff between the provision of machine-generated predictions to inform decisions and the desire to preserve managerial effort. Partial disclosure of machine information can be optimal because it balances this tradeoff. Our results establish that the optimal disclosure policy varies non-monotonically in the product's profit margin; that is, disclosure can transition from full to partial and back to full again as the margin decreases (ϕ increases). In particular, observe from Figure 3 that the policy transitions from (F) to (D) to (F) to (E) and back to (F) as ϕ increases. For very high-margin or very low-margin products, managerial incentives align with the firm's profit objective, making full disclosure optimal because it does not dampen effort. As the margin shrinks from very high to high (grows from very low to low), downplaying high predictions for high-margin products (exaggerating low predictions for low-margin products) helps retain managerial effort. When the margin is medium, however, the informational value of a fully-disclosed machine prediction (versus an obfuscated prediction) dominates the value of inducing managerial effort, making full disclosure without human effort optimal. Put differently, for medium-margin products, the firm can automate the forecasting activity and effectively disregard

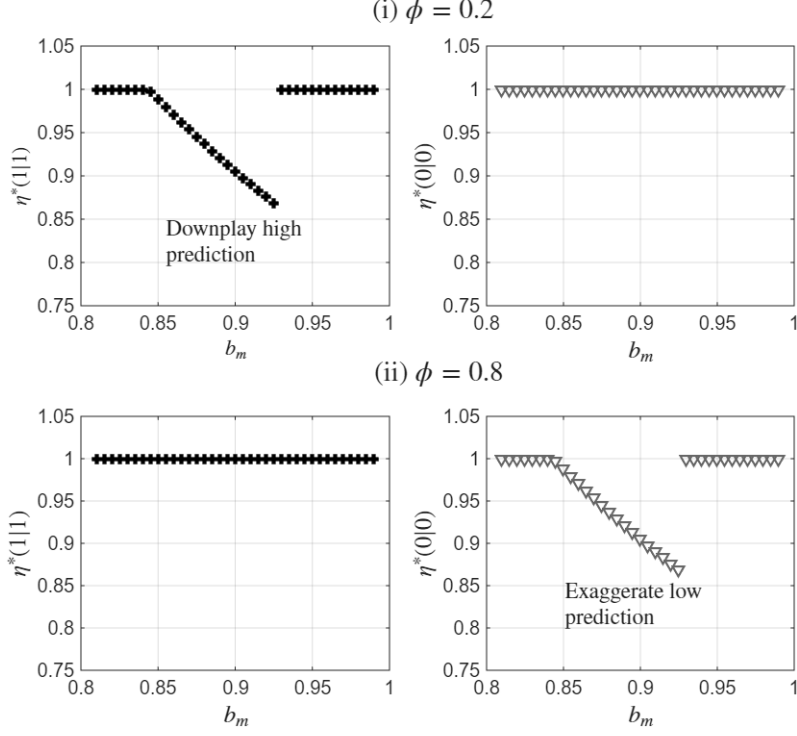


Figure 6: η_1^* and η_0^* as functions of b_m with given ϕ [Parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $k = 8.5$, $r = 10$, and $b_h = 0.8$, whereas c is varied based on ϕ]

the attendant loss of effort.

4.3.3 Effect of the Machine's Prediction Capability As discussed in Remark 2, partial disclosure can be interpreted as tuning the machine to be less informative than its prediction capability defined by b_m . Our results therefore imply that the firm may benefit from deploying a less informative machine (by introducing biases toward the high or low state depending on the product's profit margin) even if deploying the machine's full capability is costless.

Proposition 3. *All else constant, there exists a threshold $\hat{b}_m$ such that for $b_m < \hat{b}_m$ it can be optimal to obfuscate (downplay or exaggerate the machine's prediction) with a greater probability as b_m increases. In particular, (i) for $\phi \leq u_0$, the optimal policy changes from (F) to (D) and then back to (F) as b_m increases, with η_1^* starting at one when (F) optimal, then decreasing while (D) optimal, and finally jumping back to one when (F) again optimal. (ii) for $\phi \geq l_1$, the optimal policy changes from (F) to (E) and then back to (F) as b_m increases, with η_0^* starting at one when (F) optimal, then decreasing while (D) optimal, and finally jumping back to one when (F) again optimal.*

In other words, it can be optimal to garble the machine's prediction with even greater probability as the machine's prediction capability increases (unless its capability exceeds a certain threshold).

Proposition 3 is illustrated in Figure 6 where we set the human accuracy as $b_h = 0.8$ and consider two different cost-to-revenue ratios: (i) 0.2 and (ii) 0.8. For both (i) and (ii), when b_m is slightly greater than b_h , full disclosure is optimal because it does not dampen effort. However, as b_m increases further, η_1^* (case (i)) or η_0^* (case (ii)) starts to decrease in b_m because more obfuscation is required to retain managerial effort. Once b_m reaches 0.925 approximately, implementing full disclosure (deploying the machine’s full capability) again becomes optimal, but this time without managerial effort.

We emphasize that a decrease in η_1^* or η_0^* reflects a higher probability of garbling the machine prediction X_m , but that does not necessarily imply a reduction in the informativeness of predicting the market state θ . In fact, it can be shown that as the machine’s prediction capability b_m increases, the optimally-tuned informativeness in predicting the true market state θ will not decline.¹⁰

5 Extensions

5.1 Optimizing Compensation with a Full-Disclosure Machine

As discussed earlier in Remark 1, our base model is equivalent to one in which the manager is compensated through a linear performance contract: the manager receives $w_B + w_P\pi$, where w_B is the base salary, w_P is the commission rate, and π is the firm’s realized monetary profit not including compensation. With that perspective, the effort cost k in the base model reflected the true effort cost κ scaled relative to the commission rate w_P , i.e., $k = \kappa/w_P$. In effect, our analysis so far has focused on optimizing information provision so as to address managerial effort concerns while taking the compensation contract as given and unchanging. We now examine an alternative approach, compensation design, to address the manager’s incentive issue by allowing the firm to optimize the contract parameters $\{w_B, w_P\}$ when providing a full-disclosure machine.

To set a fair basis for comparison between the two different approaches, we first consider an existing (status quo) setting in which no prediction machine is available but the firm optimally designs the manager’s compensation parameters. We regard that optimal design as the *existing* contract. Then, assuming that a prediction machine becomes available, the firm may (1) retain

¹⁰As discussed in Remark 2, for predicting θ , we allow the firm to choose, among all information structures $\tilde{\eta}(s_m|\theta)$ that are (weakly) less informative than the one governed by b_m . Given any $b_m^+ > b_m$, any information structure $\tilde{\eta}(s_m|\theta)$ obtained under capability b_m can also be obtained (through garbling) under capability b_m^+ . Using this property along with the characterization of optimal disclosure policies, we can show that the optimal information structure under capability b_m^+ cannot be less informative than that under capability b_m . However, it should be noted that the optimal information structures under b_m^+ and b_m are not necessarily Blackwell comparable, and so we cannot conclude that the informativeness of market-state predictions is increasing or weakly increasing in b_m .

the existing contract while optimally disclosing the machine information as characterized in §4.3, or (2) deploy a full-disclosure machine but re-optimize the contract parameters by considering the manager's effort choice in response to machine predictions. We refer to the former strategy as the information design approach and the latter as the optimal contracting approach. In the contracting approach, we assume that the firm bears no administrative or personnel costs when altering compensation terms, such as the costs of bargaining with unions or the potential loss of long-term employees (Sandvik et al., 2021; Young and Morris, 2023).

In Appendix D, we formulate (and solve) the contract design problem in which the firm chooses $\{w_B, w_P\}$ to maximize its expected payoff $V = (1 - w_P)\mathbb{E}[\pi] - w_B$, both for the setting without any machine (status quo) and the setting with a full-disclosure machine. In both settings, the firm's choice of contract parameters $\{w_B, w_P\}$ must satisfy the incentive compatibility (IC) constraint that induces effort and the limited liability (LL) constraint that guarantees the manager a given minimum salary $\underline{t}$.¹¹ We focus on the exogenous parameter range where inducing human effort is indeed optimal under the contracting approach.¹² For the status-quo (no machine setting), we denote the optimal contract as $\{w_B^N, w_P^N\}$ and we derive this in closed form in Appendix D. For the full-disclosure machine setting, because we are interested in comparing the contracting approach to the information design approach, we focus on the exogenous parameter space for which partial disclosure would be optimal under the status-quo optimal contract; that is, the parameter space for which partial disclosure policies (D) or (E) is optimal under $\{w_B^N, w_P^N\}$. We derive (again in closed form) the optimal contract parameters for the full-information machine setting and denote them as $\{w_B^F, w_P^F\}$. We establish that in both settings, the optimal base salary terms (w_B^F and w_B^N) are set such that the LL constraint is binding. We also prove that $w_P^F > w_P^N$. In other words, to induce effort with a full-disclosure machine, the firm must raise the commission rate above its existing level, thus resulting in a higher agency cost.

Let π_{info} and $\pi_{contract}$ denote the firm's expected monetary profit (ignoring manager compensation) under the information design and optimal contracting approaches, respectively, and let V_{info} and $V_{contract}$ denote the expected overall payoffs of the firm (profit net of compensation). It is not immediately obvious whether V_{info} is less than or greater than $V_{contract}$. On the one hand, the

¹¹In practice, firms are typically required to ensure a minimum salary, regardless of realized profits. The LL constraint is common in moral hazard settings to satisfy the minimum wage requirement, which is typically determined by labor regulations or commonly accepted market standards.

¹²If the manager's effort cost κ is prohibitively high, or if the optimal order quantity is independent of X_h (which happens if $\phi \leq \frac{\lambda(1-b_h)}{1-\lambda}$ or $\phi \geq b_h$), then modifying the existing contract to induce effort will not improve the firm's expected payoff. We rule out these cases in the analysis.

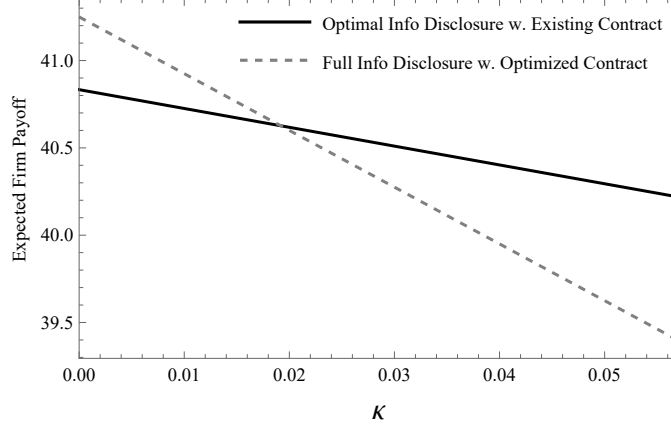


Figure 7: Optimal (partial) information disclosure without altering existing compensation terms versus full information disclosure with optimized compensation terms. (Parameter value: $\lambda = \frac{1}{2}$, $b_m = 0.9$, $b_h = 0.85$, $a_H = 20$, $a_L = 10$, $r = 10$, $c = 7$, $\underline{t} = 0$)

optimal contracting approach delivers a higher monetary profit, i.e., $\pi_{contract} > \pi_{info}$, because the manager exerts the desired effort level and has full machine information. On the other hand, the optimal contracting approach requires a higher commission rate be paid to the manager ($w_P^F > w_P^N$). We analytically establish below that (assuming equal priors) the information design approach outperforms the optimal contract approach (i.e., leads to a higher firm payoff) when the manager's original effort cost κ is large.

Proposition 4. *Assume the prior $\lambda = \frac{1}{2}$. When partial disclosure (i.e., policy D or E) is optimal under the information design approach, $V_{info} > V_{contract}$ if κ exceeds a threshold $\hat{\kappa}$.*

As demonstrated in Figure 7, the firm's expected payoff decreases in the effort cost κ . In the information design approach, the payoff decreases because the firm has to sacrifice more of the machine's predictive-value to induce managerial effort. In the contracting approach, the firm suffers from a higher agency cost (higher optimal commission rate) as effort cost increases. The firm's payoff is more sensitive to effort cost under the contracting approach, and so the information design approach is preferable (i.e., results in a higher firm payoff) once κ is large enough. Our results establish that, even if compensation adjustment is costless, firms may benefit more from optimally tuning machine information disclosure (and not changing compensation) than from fully disclosing the machine's information and re-optimizing compensation. We acknowledge that our analysis focuses on linear compensation schemes due to their widespread use in practice and in the literature. Incorporating more complex nonlinear compensation structures might be a valuable direction for future research.

5.2 Machine Less Accurate than Human (The Case of $b_m \leq b_h$)

Our analysis has so far assumed that the machine's predictive capability is superior to the human's, i.e., $b_m > b_h$. Empirical evidence suggests, however, that in some settings AI models do not consistently outperform human experts in forecasting accuracy (Abolghasemi et al., 2024). For this reason, we now consider the case of $b_m \leq b_h$.¹³

For the base model, i.e., $b_m > b_h$, we established that when partial disclosure is optimal, the firm should garble to the minimum extent necessary to induce managerial effort; see earlier disclosure policies (D) and (E) and their related discussion in §4.3.2. We will establish below that this minimum garbling result does not necessarily hold for $b_m \leq b_h$. Intuitively, when human prediction is superior to machine prediction, inducing managerial effort is highly desirable even if it comes at the cost of higher machine-information loss through more garbling. In addition to policies (D) and (E), a few more forms of partial disclosure policies can emerge in the optimal solution. It can be optimal to downplay high predictions or exaggerate low predictions with greater intensity, as defined in policies (D') and (E'):

$$\eta_1^* = \frac{\lambda - \psi_2}{1 - \psi_2} \frac{1}{\lambda} \text{ and } \eta_0^* = 1 \quad (\text{D}')$$

$$\eta_1^* = 1 \text{ and } \eta_0^* = \left(1 - \frac{\lambda}{\psi_1}\right) \frac{1}{1 - \lambda} \quad (\text{E}')$$

Note that η_1^* is smaller in (D') than in (D), that is, high predictions are downplayed to low with a larger probability. Similarly, η_0^* is smaller in (E') than in (E), that is, low predictions are exaggerated with a greater probability. Moreover, it can be optimal to simultaneously downplay high predictions and exaggerate low predictions by setting both η_1^* and η_0^* strictly less than one, as defined in policy (D&E):

$$\eta_1^* = \frac{\psi_2}{\lambda} \frac{\lambda - \psi_1}{\psi_2 - \psi_1} \text{ and } \eta_0^* = \frac{1 - \psi_1}{1 - \lambda} \frac{\psi_2 - \lambda}{\psi_2 - \psi_1} \quad (\text{D\&E})$$

We fully characterize the optimal messaging policy in Appendix E and the solution structure is illustrated in Figure 8. Below, we summarize how the optimal policies differ from the base model.

As was the case for the machine-more-accurate setting (Figure 3), full disclosure (F) is again optimal for very high-margin or very low-margin products because it does not dampen managerial effort. Also, as the margin shrinks from very high to high (grows from very low to low), downplaying

¹³Even when the machine is less accurate than the human, the machine prediction can still add value if the machine and the human's information sources are not completely overlapping; that is, the machine has access to some data that the human does not (or that the human cannot readily process).

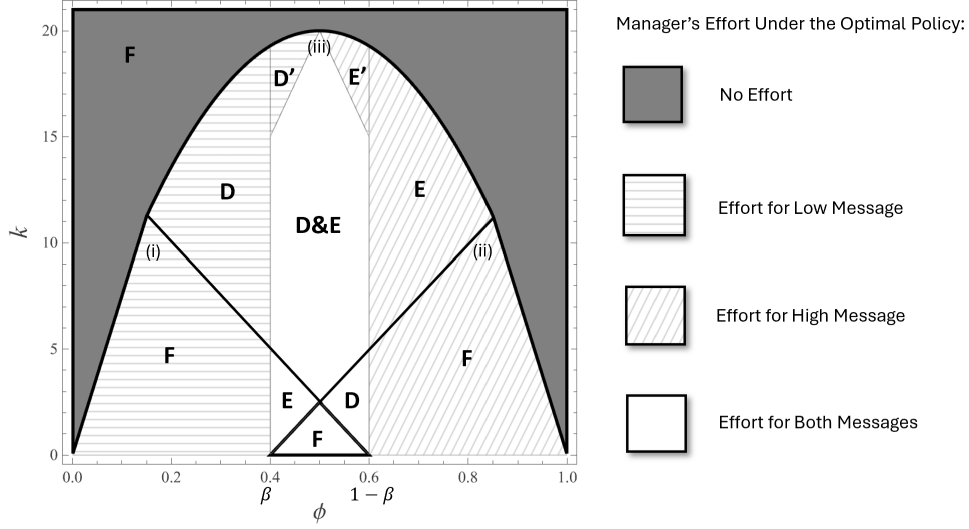


Figure 8: Structure of the optimal information provision for the case of $b_m \leq b_h$ (Parameter values: $b_m = 0.85$, $b_h = 0.9$, $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$, and c is varied based on ϕ .)

(D) high predictions for high-margin products (exaggerating (E) low predictions for low-margin products) remains optimal. For medium-margin products, the optimal disclosure policy diverges significantly from the machine-more-accurate setting where full disclosure (and foregoing managerial effort) was optimal. In the human-more-accurate setting, partial disclosure to induce effort (possibly under both message realizations) is often optimal.

In particular, consider the region $\beta < \phi \leq 1 - \beta$.¹⁴ As illustrated by the white region in Figure 8, the optimal policy varies with effort cost k . For sufficiently small k , full disclosure is optimal because the manager always has incentive to exert effort. Adjacent to this, (E) and (D) are optimal: in (E), the manager already has incentive to exert effort upon receiving a low prediction, and the optimal policy probabilistically exaggerates low predictions as high to further preserve effort under high messages; symmetrically, in (D), the manager has incentive to exert effort upon receiving a high prediction, and the optimal policy further induces effort under low messages by probabilistically downplaying high predictions. As k increases such that the manager has no incentive to exert effort under either machine prediction, garbling both prediction realizations, policy (D&E), becomes optimal.¹⁵ At even higher effort costs, (D') or (E') becomes optimal, inducing effort more often by

¹⁴The region $\beta < \phi \leq 1 - \beta$ exists only when $b_m < b_h$, since $\beta \leq \frac{1}{2}$ iff $b_m \leq b_h$.

¹⁵In our base model, when (D&E) is optimal, multiple optimal solutions exist and providing no machine information is also optimal leading to the same expected profit. This multiplicity of optimal solutions stems from the assumptions of binary effort choice and two-point distributed demand (which make the monetary profit $\Pi(v)$ piece-wise linear in v). We present the solution as (D&E) not only because we follow a tie-breaking rule that selects the most informative messaging policy, but more importantly because (D&E) is the unique optimal solution when either of those assumptions is relaxed. More details are discussed in Appendices F and G.

downplaying or exaggerating with greater intensity than in (D) or (E).¹⁶

In summary, when the machine has lower prediction capability than the human, it is always optimal to keep the human in the loop. For moderate-margin products, it is optimal to induce more effort through obfuscation rather than adopt full disclosure without engaging human effort.

5.3 General Demand Distributions

Our base model assumed demand followed a two-point distribution. We now show that our results extend to more general demand distributions. Suppose that, under each market state $\theta \in \{H, L\}$, demand d_θ follows distribution F_θ , where F_H first-order stochastically dominates F_L . For any posterior μ , the expected monetary profit under the optimal order quantity is $\pi^*(\mu) = \max_{Q \geq 0} r\mathbb{E}[\mu \min(d_H, Q) + (1 - \mu) \min(d_L, Q)] - cQ$. As in the base model, we can replace μ with the transformed posterior v associated with the machine prediction X_m . We can analytically generalize our main results when d_θ is uniformly distributed (for $\theta \in \{H, L\}$) under the following assumptions:

- (i) d_θ is either (a) uniformly distributed over $[a_\theta - \sigma, a_\theta + \sigma]$ with $a_H - a_L \geq 2\sigma$, or (b) uniformly distributed over $[0, a_\theta]$ with $a_H > a_L$.
- (ii) The unconditional predictions X_m and X_h are uncorrelated, i.e., $\rho(X_m, X_h) = 0$, or equivalently, $b_m + b_h = 1 + \lambda$.

Assumption (i) captures two common forms of state-dependent demand uncertainty when the states differ in mean demand. In (a), the two market states are non-overlapping but have the same variance; in (b) the market states are nested but have the same coefficient of variation. Under Assumptions (i) and (ii), we can prove that the value of acquiring the human prediction is unimodal in v . Consequently, as in the base model, if the effort cost k is below a threshold, there exists an interval $[\psi_1, \psi_2]$ such that $e^*(v) = 1$ iff $v \in [\psi_1, \psi_2]$. Different to the base model, closed-form expressions for ψ_1 and ψ_2 are generally unavailable. Nevertheless, the monetary profit $\Pi(v)$ is piecewise convex in v and discontinuous at $v = \psi_1$ and ψ_2 . Consequently, partial disclosure can be optimal, taking on exactly one of the forms introduced earlier. The conditions for each disclosure policy to be optimal are summarized in Proposition 5, although these conditions are more implicit than those in the base model.¹⁷

¹⁶In fact, (D) and (E) use the least garbled information that motivates effort, whereas (D') and (E') use the maximally garbled information that motivates effort.

¹⁷State-dependent uniform demand has also been adopted in the recent information design literature, e.g., Candogan and Gurkan (2024). Our setting differs substantially from their supplier-retailer setting; for example, there is no

Proposition 5. *Provided Assumptions (i)(ii) and k is lower than a threshold I_M , there exists an interval $[\psi_1, \psi_2]$ such that $e^*(v) = 1$ iff $\psi_1 \leq v \leq \psi_2$. If $\Pi(\psi_i) - \Pi(0) \leq \psi_i(\Pi(1) - \Pi(0))$ for both $i = 1, 2$, then full information disclosure is optimal. Otherwise, partial information is optimal as characterized below. (i) If $\Pi(\psi_1) - \Pi(0) > \psi_1(\Pi(1) - \Pi(0))$ and $(1 - \psi_2)(\Pi(1) - \Pi(\psi_1)) < (1 - \psi_1)(\Pi(1) - \Pi(\psi_2))$, then η^* is given by (E') for $\lambda \leq \psi_1$, and by (D) for $\lambda > \psi_1$; (ii) if $\Pi(\psi_2) - \Pi(0) > \psi_2(\Pi(1) - \Pi(0))$ and $\psi_1(\Pi(\psi_2) - \Pi(0)) > \psi_2(\Pi(\psi_1) - \Pi(0))$, then η^* is given by (E) for $\lambda \leq \psi_2$, and by (D') for $\lambda > \psi_2$; (iii) otherwise, η^* is given by (E') for $\lambda \leq \psi_1$, by (D&E) for $\psi_1 < \lambda \leq \psi_2$, and (D') for $\lambda > \psi_2$.*

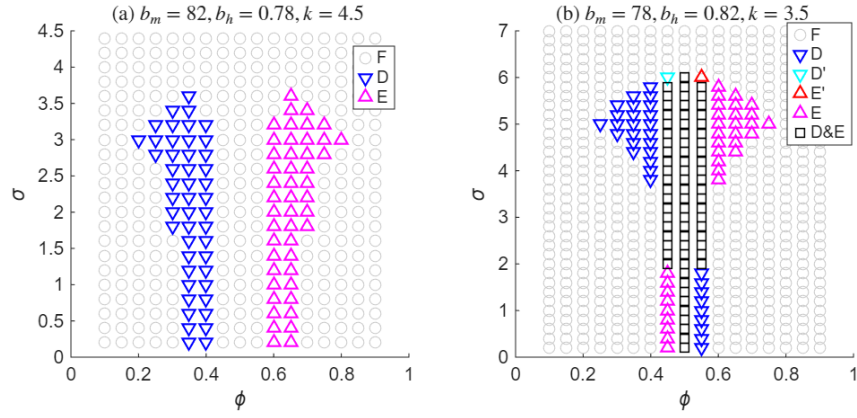


Figure 9: Effect of the standard deviation σ on the policy structure [parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$ and c is varied according to ϕ given in the graph]

We also explored non-uniform distributions by conducting a numerical study in which d_θ is normally distributed with mean a_θ and standard deviation σ . Market states thus have overlapping demand supports. Our numerical results again confirm that the optimal messaging policy has the same structure as in our base model setting with two-point distributed demand. As ϕ increases, the optimal policy changes from (F) to (D) to (F) to (E) and back to (F) when $b_m > b_h$; see Figure 9(a). Partial disclosure (in various forms) is optimal for medium value of ϕ when $b_m < b_h$; see Figure 9(b). Additionally, we observe that, for a fixed ϕ , partial disclosure is more likely to be optimal when demand has a moderate standard deviation σ . As σ increases, accurately ascertaining the market state θ becomes less important because the residual (unresolvable) demand uncertainty becomes more relevant than market-state knowledge, and so the manager has less incentive to exert effort.

effort aspect in their setting, whereas in our setting the receiver (manager) can generate an additional forecast. Consequently, our results differ. Their corollary 2 and theorem 3 show that full disclosure is always optimal for the demand specifications considered in our extension – namely, disjoint uniform intervals with equal lengths and nested intervals where one has a higher upper bound. In contrast, in our setting, partial disclosure can be optimal for both specifications.

For sufficiently large σ , full disclosure without inducing effort is optimal. Interestingly, as σ becomes smaller, the optimal policy may change from partial to full (e.g., when $\phi = 0.3$ or $\phi = 0.7$). This is because in such cases, a small enough σ , together with an appropriately high or appropriately low profit margin (low or high ϕ), already provides sufficient incentive to exert effort (for one of the machine predictions) under full disclosure, thus making partial disclosure unnecessary for the firm.

5.4 Continuous Information Elicitation Effort

In the base model, we assumed that the manager’s effort choice e was binary and that they obtained a signal $s_h = X_h$ if $e = 1$ but no information if $e = 0$. Our results can be generalized to a setting where the manager chooses a continuous effort level $e \in [0, 1]$ while incurring a convex increasing cost ke^2 where $k > 0$. Intuitively, a higher effort level should lead to a more informative prediction. To model this, we follow the “truth-or-noise” setting widely adopted in the literature (Johnson and Myatt, 2006; Hagiu and Wright, 2020). With probability e , s_h reveals the true value of X_h , i.e., $s_h = X_h$. With probability $1 - e$, s_h is an independent draw from the same distribution as X_h , and thus not informative at all. Consequently, s_h and X_h are correlated according to the following conditional probabilities: $\Pr(X_h = 1|s_h = 1) = e + (1 - e)\lambda$ and $\Pr(X_h = 0|s_h = 0) = e + (1 - e)(1 - \lambda)$. This effort-dependent prediction is also in line with Taylor and Xiao (2009). We can interpret X_h as the best prediction that a human can generate if choosing $e = 1$.

Despite the continuity of e , we can prove that the optimal effort level is either zero or sufficiently high. The intuition is as follows. Exerting effort $e > 0$ creates value only if e is large enough for the realization of s_h to meaningfully influence the order quantity. If e is too small, the posteriors conditional on $s_h = 1$ and $s_h = 0$ will not differ enough to favor different order quantities; thus, no effort is preferable to a small but inadequate effort. As a result, akin to the base model, there exists an interval $[\psi_1, \psi_2]$ such that a positive effort level is optimal iff $v \in [\psi_1, \psi_2]$ whereas no effort is optimal if v is outside $[\psi_1, \psi_2]$. Thus, partial disclosure policies work through the same mechanism as in our base model, motivating effort by inducing posterior beliefs $v = \psi_1$ or $v = \psi_2$. More detailed discussions can be found in Appendix G.

6 Conclusion

AI has many obvious benefits for business operations, but practitioners and researchers have voiced the importance of accounting for human behavior during AI deployment. Our work establishes

that AI induced effort-dampening is a valid concern in a sales and operations (demand-forecasting and procurement-quantity) context where a firm deploys a prediction machine to augment a human manager. We fully characterize the conditions under which a manager reduces forecasting effort in response to the firm deploying a full-disclosure machine. This effort withdrawal reduces the potential value of the machine, and we therefore examine whether the firm can counteract effort dampening by configuring the machine’s communication (disclosure policy) to the manager. We also compare this informational approach to a contracting approach that uses compensation design to mitigate effort withdrawal.

The optimal machine disclosure strategy depends critically on the product’s cost-to-revenue ratio. In particular, at medium ratios the firm should use full disclosure and essentially remove the human from the forecasting loop; at very high or very low ratios the firm can use full disclosure and still retain managerial effort; but at high or low ratios the firm should only partially disclose the machine prediction – either exaggerating or downplaying its prediction – to preserve managerial effort. The firm should engage in (weakly) more obfuscation as the machine’s maximum capability improves but the resulting informativeness of the optimally tuned market-state prediction does not decline as the machine’s maximum capability improves. When compared to an alternative contracting approach to manage effort dampening, the firm should implement an information-design approach unless they perceive managerial effort cost to be low, in which case compensation design can outperform information design.

Future research in this area could build on this work in a number of ways. Experimental research could study prediction-machine disclosure policies with human subjects in a newsvendor lab setting. Theoretical research could explore disclosure policies in other settings (e.g., medical diagnoses, predictive maintenance, supply-disruption adjustment) where prediction machines augment human experts.

References

- Abolghasemi, Mahdi, Odkhishig Ganbold, Kristian Rotaru. 2024. Humans vs. large language models: Judgmental forecasting in an era of advanced ai. *International Journal of Forecasting* .
- Ahmad, Sayed Fayaz, Heesup Han, Muhammad Mansoor Alam, Mohd Rehmat, Muhammad Irshad, Marcelo Arraño-Muñoz, Antonio Ariza-Montes, et al. 2023. Impact of artificial intelligence on human loss in decision making, laziness and safety in education. *Humanities and Social Sciences Communications* **10**(1) 1–14.
- Alexander, Veronika, Collin Blinder, Paul J Zak. 2018. Why trust an algorithm? performance, cognition, and neurophysiology. *Computers in Human Behavior* **89** 279–288.

- Alizamir, Saed, Francis de Véricourt, Shouqiang Wang. 2020. Warning against recurring risks: An information design approach. *Management Science* **66**(10) 4612–4629.
- Baiman, Stanley, Serguei Netessine, Richard Saouma. 2010. Informativeness, incentive compensation, and the choice of inventory buffer. *The Accounting Review* **85**(6) 1839–1860.
- Balakrishnan, Maya, Kris Johnson Ferreira, Jordan Tong. 2025. Human-algorithm collaboration with private information: Naïve advice-weighting behavior and mitigation. *Management Science* .
- Beer, Ruth, Anyan Qi, Ignacio Rios. 2024. Behavioral externalities of process automation. *Management Science* Forthcoming.
- Bimpikis, Kostas, Giacomo Mantegazza. 2023. Strategic release of information in platforms: Entry, competition, and welfare. *Operations Research* **71**(5) 1619–1635.
- Bimpikis, Kostas, Yiangos Papanastasiou, Wenchang Zhang. 2024. Information provision in two-sided platforms: Optimizing for supply. *Management Science* **70**(7) 4533–4547.
- Blackwell, David. 1953. Equivalent comparisons of experiments. *The annals of mathematical statistics* 265–272.
- Boyacı, Tamer, Caner Canyakmaz, Francis de Véricourt. 2024. Human and machine: The impact of machine input on decision making under cognitive limitations. *Management Science* **70**(2) 1258–1275.
- Boyaci, Tamer, Souptika Chakraborty, Huseyin Gurkan. 2024. Persuading skeptics and fans in the presence of additional information. *Production and Operations Management* **33**(5) 1142–1154.
- Candogan, Ozan. 2020. Information design in operations. *Pushing the Boundaries: Frontiers in Impactful OR/OM Research*. INFORMS, 176–201.
- Candogan, Ozan, Huseyin Gurkan. 2024. The value of information design in supply chain management. *Management Science* Forthcoming.
- Chen, Yonghui, Ailing Xu, Qiao-Chu He, Ying-Ju Chen. 2023. Smart navigation via strategic communications in a mixed autonomous paradigm. *Production and Operations Management* 1–17DOI: 10.1111/poms.13987.
- Chu, Leon Yang, Noam Shamir, Hyoduk Shin. 2017. Strategic communication for capacity alignment with pricing in a supply chain. *Management Science* **63**(12) 4366–4388.
- Cui, Ruomeng, Meng Li, Shichen Zhang. 2022. Ai and procurement. *Manufacturing & Service Operations Management* **24**(2) 691–706.
- Dai, Tinglong, Shubhranshu Singh. 2024. Using ai as gatekeeper or second opinion: Designing patient pathways for ai-augmented healthcare. *Available at SSRN* .
- Dai, Tinglong, Shubhranshu Singh. 2025. Artificial intelligence on call: The physician’s decision of whether to use ai in clinical practice. *Journal of Marketing Research* Forthcoming.
- Dai, Tinglong, Jayashankar M Swaminathan. 2025. Artificial intelligence and operations: A foundational framework of emerging research and practice. *Production and Operations Management* 10591478251412943.
- Dai, Tinglong, Terry Taylor. 2025. Designing enterprise ai systems: Hallucination, creativity, and moral hazard. *Creativity, and Moral Hazard (December 31, 2025)* Available at SSRN <https://ssrn.com/abstract=5996714>.
- Dai, Tinglong, Sridhar Tayur. 2022. Designing ai-augmented healthcare delivery systems for physician buy-in and patient acceptance. *Production and Operations Management* **31**(12) 4443–4451.
- de Véricourt, Francis, Huseyin Gurkan. 2023. Is your machine better than you? you may never know. *Management Science* .

- De Véricourt, Francis, Huseyin Gurkan, Shouqiang Wang. 2021. Informing the public about a pandemic. *Management Science* **67**(10) 6350–6357.
- Dell’Acqua, Fabrizio. 2022. Falling asleep at the wheel: Human/ai collaboration in a field experiment on hr recruiters. *Work. Pap.*
- Drakopoulos, Kimon, Shobhit Jain, Ramandeep Randhawa. 2021. Persuading customers to buy early: The value of personalized information provisioning. *Management Science* **67**(2) 828–853.
- Farahani, Mehdi H, Milind Dawande, Ganesh Janakiraman, Shouqiang Wang. 2024. Avoiding fields on fire: Information dissemination policies for environmentally safe crop-residue management. *Management Science*.
- Guan, Xiaotong, Anyan Qi, Shouqiang Wang. 2025. Why the best machine may not be the best: Incentivizing human-machine collaboration. Available at <https://ssrn.com/abstract=5283571>.
- Hagiu, Andrei, Julian Wright. 2020. Platforms and the exploration of new products. *Management Science* **66**(4) 1527–1543.
- Hou, Ting, Meng Li, Yinliang Tan, Huazhong Zhao. 2024. Physician adoption of ai assistant. *Manufacturing & Service Operations Management* **26**(5) 1639–1655.
- Ibrahim, Rouba, Song-Hee Kim, Jordan Tong. 2021. Eliciting human judgment for prediction algorithms. *Management Science* **67**(4) 2314–2325.
- Iyer, Ganesh, Chakravarthi Narasimhan, Rakesh Niraj. 2007. Information and inventory in distribution channels. *Management Science* **53**(10) 1551–1561.
- Jiang, Baojun, Lin Tian, Yifan Xu, Fuqiang Zhang. 2016. To share or not to share: Demand forecast sharing in a distribution channel. *Marketing Science* **35**(5) 800–809.
- Jin, Li. 2002. CEO compensation, diversification, and incentives. *Journal of financial Economics* **66**(1) 29–63.
- Johnson, Justin P, David P Myatt. 2006. On the simple economics of advertising, marketing, and product design. *American Economic Review* **96**(3) 756–784.
- Kamenica, Emir, Matthew Gentzkow. 2011. Bayesian persuasion. *American Economic Review* **101**(6) 2590–2615.
- Klingbeil, Artur, Cassandra Grützner, Philipp Schreck. 2024. Trust and reliance on ai—an experimental study on the extent and costs of overreliance on ai. *Computers in Human Behavior* **160** 108352.
- Li, Meng, Tao Li. 2022. Ai automation and retailer regret in supply chains. *Production and Operations Management* **31**(1) 83–97.
- Lipnowski, Elliot, Laurent Mathevet, Dong Wei. 2020. Attention management. *American Economic Review: Insights* **2**(1) 17–32.
- Long, Xiaoyang, Javad Nasiry. 2020. Wage transparency and social comparison in sales force compensation. *Management Science* **66**(11) 5290–5315.
- Lu, Tao, Brian Tomlin. 2024. Credible cheap-talk communication of private demand information on both the forecast average and accuracy. *Management Science* Forthcoming.
- Matyskova, Ludmila, Alfonso Montes. 2023. Bayesian persuasion with costly information acquisition. *Journal of Economic Theory* **211** 105678.
- Nair, Devadrita, Maria Jesus Saenz. 2024. Pair people and ai for better product demand forecasting. *MIT Sloan Management Review* **65**(2) 1–5.
- Ni, Xiao, Yiwei Wang, Tianjun Feng, Lauren Xiaoyuan Lu, Yitong Wang, Congyi Zhou. 2024. Generative ai in

- action: Field experimental evidence on worker performance in e-commerce customer service operations. *Action: Field Experimental Evidence on Worker Performance in E-Commerce Customer Service Operations (November 06, 2024)* .
- Özer, Özalp, Yanchong Zheng, Kay-Yut Chen. 2011. Trust in forecast information sharing. *Management Science* **57**(6) 1111–1137.
- Papanastasiou, Yiangos, Kostas Bimpikis, Nicos Savva. 2018. Crowdsourcing exploration. *Management Science* **64**(4) 1727–1746.
- Revilla, Elena, Maria Jesus Saenz, Matthias Seifert, Ye Ma. 2023. Human–artificial intelligence collaboration in prediction: A field experiment in the retail industry. *Journal of Management Information Systems* **40**(4) 1071–1098.
- Sandvik, Jason, Richard Saouma, Nathan Seegert, Christopher Stanton. 2021. Employee responses to compensation changes: Evidence from a sales firm. *Management Science* **67**(12) 7687–7707.
- Siemens, Enno, Eirini Spiliotopoulou. 2024. Enno siemens [pdf] from ssrn.com chicken or the egg? forecast and order adjustments in supply-chain planning Available at SSRN.
- Tang, Wenjie, Karan Girotra. 2017. Using advance purchase discount contracts under uncertain information acquisition cost. *Production and Operations Management* **26**(8) 1553–1567.
- Taylor, Terry A, Wenqiang Xiao. 2009. Incentives for retailer forecasting: Rebates vs. returns. *Management Science* **55**(10) 1654–1669.
- Young, E., H. Morris. 2023. AI at work: Safety and nlra best practices for employers Available at <https://www.proskauer.com/pub/ai-at-work-safety-and-nlra-best-practices-for-employers> accessed Oct 21, 2025.

Online Appendices for “Augmenting the Operations Manager with a Prediction Machine”

A Transforming the Expected Monetary Profit as a Function of v

In Lemma A.1 below, we establish a mapping from v , the posterior belief that $\Pr(X_m = 1|s_m)$, to $\mu(v)$, the posterior belief that $\theta = H$ given v . Specifically, for each possible information set $\mathbf{s} = (s_m, \emptyset)$, $(s_m, 1)$ and $(s_m, 0)$, we can write the posterior probability $\Pr(\theta = H|\mathbf{s})$ as $\mu_N(v)$, $\mu(v|1)$ and $\mu(v|0)$, respectively. Note that $\mu_N(v)$ represents the manager’s belief regarding the market state based solely on the machine, and $\mu(v|s_h)$ denotes the posterior belief when the manager acquires an additional signal s_h .

Lemma A.1. *Suppose that, upon receiving any given machine signal $s_m \in \{0, 1\}$, the manager has a posterior belief $v = \Pr(X_m = 1|s_m)$. Then, for any given v , the following holds true.*

- *If the manager does not exert effort to generate a prediction s_h , then their posterior belief that the market state is high, $\Pr(\theta = H|s_m, \emptyset)$, can be expressed as*

$$\mu_N(v) = \frac{v(b_m - \lambda) + \lambda(1 - b_m)}{1 - \lambda}. \quad (\text{A.1})$$

- *If the manager exerts effort to generate a prediction s_h , then*

- *$s_h = 1$ with probability $\Pr(s_h = 1|s_m) = z(v) := \frac{v(b_m + b_h - 1 - \lambda) + \lambda(2 - b_m - b_h)}{1 - \lambda}$, and the resulting posterior belief $\Pr(\theta = H|s_m, 1)$ can be expressed as*

$$\mu(v|1) = \frac{v(b_m + b_h - 1 - b_h\lambda) + \lambda(1 - b_m)}{v(b_m + b_h - 1 - \lambda) + \lambda(2 - b_m - b_h)}; \quad (\text{A.2})$$

- *$s_h = 0$ with probability $\Pr(s_h = 0|s_m) = 1 - z(v)$, and the resulting posterior belief $\Pr(\theta = H|s_m, 0)$ can be expressed as*

$$\mu(v|0) = \frac{v(1 - b_h)(1 - \lambda)}{(1 - \lambda) - v(b_m + b_h - 1 - \lambda) - \lambda(2 - b_m - b_h)}. \quad (\text{A.3})$$

Using the expressions in Lemma A.1 and the profit function $\pi^*(\mu)$ derived in Lemma 1, we can rewrite the expected monetary profit given any information set $\mathbf{s}$ as a function of the posterior belief v . Specifically, we write $\Pi(s_m, \emptyset)$ as $\Pi_N(v) = \pi^*(\mu_N(v))$, representing the expected profit

without generating a human signal s_h ; and we write $\Pi(s_m, s_h)$ as $\pi^*(\mu(v|s_h))$ for each $s_h = 1, 0$, representing the expected profits conditional on $s_h = 1$ and $s_h = 0$, respectively. Define $\Pi_I(v) = \sum_{s_h \in \{0,1\}} \Pr(s_h|s_m) \pi^*(\mu(v|s_h))$ as the expected monetary profit when generating a prediction (i.e., eliciting signal s_h). Consequently, and with a slight abuse of notation, the optimal prediction effort can be recast as a function of the posterior v , denoted by $e^*(v)$, and we have $e^*(v) = 1$ iff $\Pi_I(v) - \Pi_N(v) \geq k$.

Proof of Lemma A.1. In our proofs, it is sometimes more convenient to work with the probability distribution of X_h conditional on X_m . For $i, j \in \{0, 1\}$, we define $y_{ij} = \Pr(X_h = j | X_m = i) = \frac{p_{ij}}{\sum_{j \in \{0,1\}} p_{ij}}$ where the p_{ij} 's are the joint probability distribution of (X_m, X_h) given in Table 1. Mapping to our original notation, we have

$$y_{11} = b_m + b_h - 1, y_{10} = 2 - b_m - b_h, y_{01} = \frac{(2 - b_m - b_h)\lambda}{1 - \lambda} \text{ and } y_{00} = \frac{1 - (3 - b_m - b_h)\lambda}{1 - \lambda}. \quad (\text{A.4})$$

By the assumption that X_m and X_h are either independent or positively correlated (or equivalently, $b_m + b_h \geq 1 + \lambda$), we have $y_{11} + y_{00} \geq 1$ where $y_{11} + y_{00} = 1$ represents the case where X_m and X_h are independent.

We start by deriving a general result that applies to any effort level $e \in [0, 1]$ according to the setting in Appendix G. Specifically, with probability e , s_h reveals the true value of X_h , i.e., $s_h = X_h$. With probability $1 - e$, s_h is an independent draw per the same distribution of X_h . As a result, $\Pr(X_h = 1 | s_h = 1) = e + (1 - e)\lambda$ and $\Pr(X_h = 0 | s_h = 0) = e + (1 - e)(1 - \lambda)$. Letting $e = 0$ and $e = 1$ will yield the result for the base model where $e \in \{0, 1\}$.

Given any posterior belief v regarding $X_m = 1$ and any effort level e , we derive the human's posterior belief regarding the market state θ upon receiving any realized s_h . Let $\mu(e, v | s_h)$ denote the probability conditional on s_h , which depends on both v and e .

Upon receiving $s_h = 1$, the posterior belief $\mu(e, v | 1) = \Pr(\theta = H | s_h = 1)$ can be derived as follows. Let $\Omega = \{0, 1\}$. By the law of total probability, we have

$$\begin{aligned} \Pr(\theta = H | s_h = 1) &= \sum_{\omega, \omega' \in \Omega} \Pr(\theta = H | X_m = \omega, X_h = \omega') \Pr(X_m = \omega, X_h = \omega' | s_h = 1) \\ &= \Pr(X_m = 1, X_h = 1 | s_h = 1) + \beta \Pr(X_m = 1, X_h = 0 | s_h = 1) \\ &\quad + (1 - \beta) \Pr(X_m = 0, X_h = 1 | s_h = 1) \end{aligned}$$

where the second equality follows by invoking the values of $\Pr(\theta = H | X_m, X_h)$ in Table 1. Thus, it suffices to derive probabilities $\Pr(X_m = \omega, X_h = \omega' | s_h = 1)$ for $\omega, \omega' \in \Omega$. By Bayes' rule, it follows

that

$$\begin{aligned}
& \Pr(X_m = 1, X_h = 1 | s_h = 1) \\
&= \frac{\Pr(s_h = 1 | X_m = 1, X_h = 1) \Pr(X_h = 1 | X_m = 1) \Pr(X_m = 1)}{\sum_{\omega, \omega' \in \Omega} \Pr(s_h = 1 | X_m = \omega, X_h = \omega') \Pr(X_h = \omega' | X_m = \omega) \Pr(X_m = \omega)} \\
&= \frac{(e + (1 - e)\lambda)y_{11}v}{(e + (1 - e)\lambda)y_{11}v + (1 - e)\lambda y_{10}v + (e + (1 - e)\lambda)y_{01}(1 - v) + (1 - e)\lambda y_{00}(1 - v)}
\end{aligned} \tag{A.5}$$

where we note that, for any $\omega \in \{0, 1\}$, $\Pr(s_h = 1 | X_m = \omega, X_h = 1) = e + (1 - e)\lambda$ and $\Pr(s_h = 1 | X_m = \omega, X_h = 0) = (1 - e)\lambda$ because of the truth-or-noise setting of signal s_h . Similarly, we have

$$\Pr(X_m = 1, X_h = 0 | s_h = 1) = \frac{(1 - e)\lambda y_{10}v}{(e + (1 - e)\lambda)y_{11}v + (1 - e)\lambda y_{10}v + (e + (1 - e)\lambda)y_{01}(1 - v) + (1 - e)\lambda y_{00}(1 - v)}$$

and

$$\Pr(X_m = 0, X_h = 1 | s_h = 1) = \frac{(e + (1 - e)\lambda)y_{01}(1 - v)}{(e + (1 - e)\lambda)y_{11}v + (1 - e)\lambda y_{10}v + (e + (1 - e)\lambda)y_{01}(1 - v) + (1 - e)\lambda y_{00}(1 - v)}.$$

Substituting the above expressions into (A.5) yields

$$\Pr(\theta = H | s_h = 1) = \frac{(e + (1 - e)\lambda)y_{11}v + \beta(1 - e)\lambda y_{10}v + (1 - \beta)(e + (1 - e)\lambda)y_{01}(1 - v)}{(e + (1 - e)\lambda)y_{11}v + (1 - e)\lambda y_{10}v + (e + (1 - e)\lambda)y_{01}(1 - v) + (1 - e)\lambda y_{00}(1 - v)}.$$

Invoking the expressions of β , y_{11} , y_{10} and y_{01} in terms of b_m and b_h , we can simplify $\Pr(\theta = H | s_h = 1)$ as

$$\mu(v, e | 1) = \frac{e(1 - \lambda)[v(b_h + b_m - 1 - \lambda) + \lambda(1 - b_m)] + \lambda(\lambda(1 - b_m) + v(b_m - \lambda))}{\lambda(1 - \lambda) + e(b_m + b_h - 1 - \lambda)(v - \lambda)}. \tag{A.6}$$

Similarly, we can derive $\Pr(\theta = H | s_h = 0)$ as

$$\mu(v, e | 0) = \frac{(1 - \lambda)[\lambda(1 - b_m) + v(b_m - \lambda) - e(v(b_m + b_h - 1 - \lambda) + \lambda(1 - b_m))]}{(1 - \lambda)^2 - e(b_m + b_h - 1 - \lambda)(v - \lambda)}. \tag{A.7}$$

Now letting $e = 0$, s_h is uninformative at all and, indeed, we have $\mu(v, 0 | 1) = \mu(v, 0 | 0) = \frac{v(b_m - \lambda) + \lambda(1 - b_m)}{1 - \lambda}$. This probability is written as $\mu_N(v)$ in Lemma A.1, denoting the posterior belief without eliciting s_h . Letting $e = 1$, we obtain the expressions (A.2) and (A.3) stated in Lemma A.1.

Moreover, with belief $v = \Pr(X_m = 1)$, the probability of $s_h = 1$, denoted by $z(v, e)$, is given by

$$\begin{aligned}\Pr(s_h = 1) &= \sum_{\omega, \omega' \in \Omega} \Pr(s_h = 1 | X_m = \omega, X_h = \omega') \Pr(X_h = \omega' | X_m = \omega) \Pr(X_m = \omega) \\ &= (e + (1 - e)\lambda)y_{11}v + (1 - e)\lambda y_{10}v + (e + (1 - e)\lambda)y_{01}(1 - v) + (1 - e)\lambda y_{00}(1 - v).\end{aligned}$$

Substituting (A.4) and $e = 1$, we have $z(v, 1) = \frac{v(b_m + b_h - 1 - \lambda) + \lambda(2 - b_m - b_h)}{1 - \lambda}$, which is the probability $z(v)$ in Lemma A.1. $\square$

B Proofs

B.1 Proof of Lemma 1

Proof. With belief μ , the profit function $\Pi(Q) = r[\mu \min(a_\theta, Q) + (1 - \mu) \min(a_L, Q)] - cQ$ is piecewise linear and concave in Q . So, the optimal quantity is attained either at $Q = a_L$ or $Q = a_H$; we break the tie by setting $Q = a_L$. The lemma follows by choosing the higher profit between $\pi(a_L) = (r - c)a_L$ and $\pi(a_H) = (r - c)a_L + (r\mu - c)(a_H - a_L)$. $\square$

B.2 Proof of Proposition 1

Proof. **(i) No Machine Information.** Without machine, the manager decides on e with prior be-

lief λ . If $e^* = 0$, then it obtains an expected monetary profit $\Pi_N^n = \begin{cases} (r - c)a_L & \text{if } \lambda \leq \phi \\ (r - c)a_L + (r\lambda - c)(a_H - a_L) & \text{if } \lambda > \phi. \end{cases}$

If $e^* = 1$, then by Assumption 1, it will receive a signal $s_h = 1$ and update the posterior belief in $\theta = H$ to $\mu = b_h$ with probability λ , while receiving $s_h = 0$ and updating the posterior belief to $\mu = \frac{\lambda(1 - b_h)}{1 - \lambda}$ with probability $1 - \lambda$. Note that $\frac{\lambda(1 - b_h)}{1 - \lambda} < \lambda < b_h$ for $b_h \in (\lambda, 1]$. By Lemma 1, the expected monetary profit with information elicitation is comprised of three cases:

$$\Pi_I^n = \begin{cases} (r - c)a_L & \text{if } b_h \leq \phi, \\ (r - c)a_L + \lambda(rb_h - c)(a_H - a_L) & \text{if } \frac{\lambda(1 - b_h)}{1 - \lambda} \leq \phi < b_h, \\ (r - c)a_L + (r\lambda - c)(a_H - a_L) & \text{if } \frac{\lambda(1 - b_h)}{1 - \lambda} > \phi \end{cases}$$

because the optimal order size is always a_L and a_H units, regardless of the realized s_h , for the first and third cases, respectively, whereas it is a_H units conditional on $s_h = 1$ but a_L units conditional on $s_h = 0$.

Therefore, the value of acquiring s_h is given by

$$\Pi_I^n - \Pi_N^n = \begin{cases} 0 & \text{if } b_h \leq \phi, \\ \lambda(rb_h - c)(a_H - a_L) & \text{if } \lambda \leq \phi < b_h, \\ ((1 - \lambda)c - \lambda r(1 - b_h))(a_H - a_L) & \text{if } \frac{\lambda(1 - b_h)}{1 - \lambda} \leq \phi < \lambda, \\ 0 & \text{if } \frac{\lambda(1 - b_h)}{1 - \lambda} > \phi. \end{cases}$$

Since $e^* = 1$ iff $\Pi_I - \Pi_N \geq k$, it follows that $e^* = 1$ iff $\hat{k} := \frac{k}{r(a_H - a_L)} \leq \min((1 - \lambda)\phi - \lambda(1 - b_h), \lambda(b_h - \phi))$. By rearranging the terms, this is equivalent to $l_n \leq \phi \leq u_n$.

(ii) Full Machine Information. The result for full machine information is a special case of Lemma 2, where the optimal effort is characterized for any posterior $v = \Pr(X_m = 1|s_m)$. Under full machine information, with probability λ (resp. $1 - \lambda$), $s_m = 1$ leading to posterior $v = 1$ (resp. $s_m = 0$ leading to $v = 0$). The optimal conditions for $e^*(0) = 1$ and $e^*(1) = 1$ are obtained by invoking Cases (i) and (ii) in Lemma 2. $\square$

B.3 Proof of Lemma 2

Proof. We prove the results by considering both cases of $\beta > 1/2$ and $\beta \leq 1/2$. The former case is presented in Lemma 2 of the main paper, and the general statement allowing for the latter case is in Lemma E.1 of Appendix E. We will use notation y_{ij} defined in (A.4) along with $\beta = \frac{1 - b_h}{2 - b_m - b_h}$. The optimal effort $e^*(v) = 1$ iff $\Pi_I(v) - \Pi_N(v) \geq k$. Depending on the three posterior beliefs $\mu(v)$, $\mu(v|0)$ and $\mu(v|1)$ relative to $\phi = c/r$, the expression of $\Pi_I(v) - \Pi_N(v)$ consists of four cases. Also, note that $\mu(v|0) \leq \mu(v) \leq \mu(v|1)$.

Case (1) For $\mu(v|1) \leq \phi$, the optimal order quantity $Q^*(\mu) = a_L$ for all posteriors. Thus, $\Pi_N(v) = \Pi_I(v) = (r - c)a_L$.

Case (2) For $\mu(v) \leq \phi < \mu(v|1)$, $Q^*(\mu) = a_L$ for both $\mu = \mu(v)$ and $\mu(v|0)$, but $Q^*(\mu) = a_H$ for $\mu = \mu(v|1)$. Thus, $\Pi_N(v) = (r - c)a_L$ and

$$\Pi_I(v) = (r - c)a_L + z(v)(r\mu(v|1) - c)(a_H - a_L) = A_I v + B_I \quad (\text{B.1})$$

where $z(v)$ is the probability of $s_h = 1$ derived in Lemma A.1, and

$$A_I = (a_H - a_L)((1 - y_{00})\beta r + (r - c)(y_{11} + y_{00} - 1)) \quad (\text{B.2})$$

and

$$B_I = (r - c)a_L + (a_H - a_L)(1 - y_{00})(r(1 - \beta) - c). \quad (\text{B.3})$$

Case (3) For $\mu(v|0) \leq \phi < \mu(v)$, $Q^*(\mu) = a_H$ for both $\mu = \mu(v)$ and $\mu(v|1)$, but $Q^*(\mu) = a_L$ for $\mu = \mu(v|0)$. Thus, $\Pi_N(v)$ and $\Pi_I(v)$ can be written as follows:

$$\Pi_N(v) = A_N v + B_N \quad (\text{B.4})$$

where

$$A_N = (a_H - a_L)((y_{11} + y_{00} - 1)(1 - \beta) + \beta)r \quad (\text{B.5})$$

and

$$B_N = (r - c)a_L + (a_H - a_L)(r(1 - y_{00})(1 - \beta) - c). \quad (\text{B.6})$$

$$\Pi_I(v) = (r - c)a_L + z(v)(r\mu(v|1) - c)(a_H - a_L) = A_I v + B_I \quad (\text{B.7})$$

where A_I and B_I are as defined in (B.2) and (B.3).

Case (4) For $\mu(v|0) > \phi$, $Q^*(\mu) = a_H$ for all the posteriors, and so

$$\Pi_N(v) = (r - c)a_L + (r\mu(v) - c)(a_H - a_L) \quad (\text{B.8})$$

and, since $z(v)\mu(v|1) + (1 - z(v))\mu(v|0) = \mu(v)$ (i.e., the law of iterated expectations),

$$\Pi_I(v) = (r - c)a_L + (r\mu(v) - c)(a_H - a_L). \quad (\text{B.9})$$

Therefore,

$$\Pi_N(v) = \Pi_I(v) = A_N v + B_N$$

where A_N and B_N are defined in (B.5) and (B.6).

Let $t_N(\phi)$, $t_1(\phi)$ and $t_0(\phi)$ denote the inverse functions of $\mu(v)$, $\mu(v|1)$ and $\mu(v|0)$, respectively.

Omitting their argument ϕ for ease of notation, we have

$$t_n = \frac{\phi - (1 - y_{00})(1 - \beta)}{(y_{11} + y_{00} - 1)(1 - \beta) + \beta}, \quad (\text{B.10})$$

$$t_1 = \frac{(1 - y_{00})(\phi + \beta - 1)}{(1 - \phi)(y_{11} + y_{00} - 1) + \beta(1 - y_{00})} \quad (\text{B.11})$$

and

$$t_0 = \frac{y_{00}\phi}{\phi(y_{11} + y_{00} - 1) + \beta(1 - y_{11})} \quad (\text{B.12})$$

Putting together the four cases and using the inverse functions t_n , t_1 and t_0 , we can write $H(v) = \Pi_I(v) - \Pi_N(v)$ as follows.

$$H(v) = \begin{cases} 0 & \text{if } v \leq t_1, \\ A_I v + B_I - (r - c)a_L & \text{if } t_1 < v \leq t_n, \\ (A_I - A_N)v + B_I - B_N & \text{if } t_n < v \leq t_0, \\ 0 & \text{if } v > t_0. \end{cases} \quad (\text{B.13})$$

First, it is easy to see that $H(v)$ is continuous in v . Second, by the assumption of $y_{11} + y_{00} - 1 \geq 0$, and it follows that $A_I \geq 0$. Moreover,

$$A_I - A_N = -(a_H - a_L)((1 - y_{11})\beta r + (y_{11} + y_{00} - 1)c) \leq 0.$$

Therefore, $H(v)$ is unimodal in v , maximized at $v = t_n$ and minimized to zero for $v \leq t_1$ and $v \geq t_0$. Recall that $e^*(v) = 1$ iff $H(v) \geq k$. It follows that (i) if $k > H(t_n)$, $e^*(v) = 0$ for all v ; and (ii) if $k \leq H(t_n)$, $e^*(v) = 1$ iff $\psi_1 \leq v \leq \psi_2$ where ψ_1 and ψ_2 are the intersections between $H(v)$ and k in the second and third cases of (B.13), respectively. Straightforward algebra yields the following expressions of ψ_1 and ψ_2 :

$$\psi_1 = \frac{k/(r(a_H - a_L)) - (1 - y_{00})(1 - \beta - \phi)}{(1 - y_{00})\beta + (1 - \phi)(y_{00} + y_{11} - 1)} \quad (\text{B.14})$$

and

$$\psi_2 = \frac{\phi y_{00} - k/(r(a_H - a_L))}{(1 - y_{11})\beta + \phi(y_{11} + y_{00} - 1)} \quad (\text{B.15})$$

Note that thresholds, $\psi_1 \leq \psi_2$, may or may not fall within $[0, 1]$. To characterize the solution structure, we define a set of functions of ϕ :

$$g_{10}(\phi) = -(1 - y_{00})\phi + (1 - y_{00})(1 - \beta),$$

$$g_{11}(\phi) = -y_{11}\phi + y_{11},$$

$$g_{20}(\phi) = y_{00}\phi,$$

$$g_{21}(\phi) = (1 - y_{11})\phi - (1 - y_{11})\beta,$$

and denote by g_{ij}^{-1} their inverse functions. Furthermore, let

$$g_m(\phi) = \frac{(y_{11} + y_{00} - 1)(1 - \phi)\phi + (1 - y_{11})(1 - y_{00})(1 - \beta)\beta}{(y_{11} + y_{00} - 1)(1 - \beta) + \beta}.$$

The following relations hold: $\psi_1 \leq 0 \iff \hat{k} \leq g_{10}(\phi)$, $\psi_1 \leq 1 \iff \hat{k} \leq g_{11}(\phi)$, $\psi_2 \leq 0 \iff \hat{k} \geq g_{20}(\phi)$, $\psi_2 \leq 1 \iff \hat{k} \geq g_{21}(\phi)$ and $k > H(t_n) \iff \hat{k} > g_m(\phi)$, where $\hat{k} = \frac{k}{r(a_H - a_L)}$.

Because $H(v)$ is unimodal, we have the following cases.

- **Case (N): Never Exerting Effort.** If $k > H(t_n)$ or $\psi_1 > 1$ or $\psi_2 < 0$, $e^*(v) = 0$ for all $v \in [0, 1]$. Using the functions defined above, $k > H(t_n)$ or $\psi_1 > 1$ or $\psi_2 < 0$ can be recast as $\hat{k} > \bar{K} := \min(g_m(\phi), g_{20}(\phi), g_{11}(\phi))$.

- **Case (A): Always Exerting Effort.** If $\psi_1 \leq 0 < 1 \leq \psi_2$, then $e^*(v) = 1$ for all $v \in [0, 1]$. The condition for this case can be written as $\hat{k} \leq \underline{K} := \min(g_{10}(\phi), g_{21}(\phi))$. Additionally, note that

$$\hat{k} \leq g_{10}(\phi) \iff \phi \leq \frac{(1 - \beta)(1 - y_{00}) - \hat{k}}{1 - y_{00}}$$

and

$$\hat{k} \leq g_{21}(\phi) \iff \phi \geq \frac{\hat{k} + \beta(1 - y_{11})}{1 - y_{11}}.$$

It follows that a nonnegative $\hat{k}$ exists to satisfy both inequalities above iff $\beta \leq \frac{1}{2}$. In other words, Case (A) emerges only when $\beta \leq \frac{1}{2}$.

- **Posterior-Dependent Effort.** Otherwise, $e^*(v)$ will depend on the value of v . We have three specific cases, depending on how interval $[\psi_1, \psi_2]$ overlaps with $[0, 1]$. Let $g_{10}^{-1}(\hat{k}) = 1 - \beta - \frac{\hat{k}}{1 - y_{00}}$ and $g_{21}^{-1}(\hat{k}) = \beta + \frac{\hat{k}}{1 - y_{11}}$ be the inverse function of g_{10} and g_{21} .

- If $\psi_1 \leq 0$ and $\psi_2 \leq 1$, $e^*(v) = 1$ iff $v \leq \psi_2$. The condition can be written with inverse functions: $\phi \leq \min(g_{10}^{-1}(\hat{k}), g_{21}^{-1}(\hat{k}))$.
- If $0 \leq \psi_1$ and $1 \leq \psi_2$, then $e^*(v) = 1$ iff $v \geq \psi_1$. The condition can be written as $\phi \geq \max(g_{10}^{-1}(\hat{k}), g_{21}^{-1}(\hat{k}))$.
- If $0 < \psi_1 \leq \psi_2 < 1$, then $e^*(v) = 1$ iff $\psi_1 \leq v \leq \psi_2$. The condition can be written as $g_{10}^{-1}(\hat{k}) < \phi < g_{21}^{-1}(\hat{k})$. Note that $\psi_1 \leq \psi_2$ is implied by $k \leq H(t_n(\phi))$.

To simplify the notation, let $l_0 = g_{20}^{-1}(\phi)$, $u_0 = g_{10}^{-1}(\phi)$, $l_1 = g_{21}^{-1}(\phi)$, and $u_1 = g_{11}^{-1}(\phi)$. Note that per (A.4), one can rewrite all these thresholds in terms of the original notation (λ, b_m and b_h) as given in the main paper. Furthermore, for $\beta > \frac{1}{2}$, we always have $u_0 < l_1$ for any positive $\hat{k}$. We can therefore simplify the above cases as stated in Lemma 2 for $\beta > \frac{1}{2}$, whereas the general case allowing for $\beta \leq \frac{1}{2}$ is presented in E.1.

Finally, we establish that, provided $0 < \psi_i < 1$, ψ_i is increasing in ϕ for both $i = 1, 2$.

For ψ_1 , its numerator is linear in ϕ with coefficient $(1 - y_{00})$, thus increasing in ϕ . Its denominator is linear in ϕ with coefficient $-(y_{11} + y_{00} - 1) \leq 0$ by the assumption of $y_{11} + y_{00} \geq 1$. Thus, the numerator is increasing while the denominator is decreasing in ϕ . Given that both the numerator and denominator are positive, ψ_1 is increasing in ϕ .

For ψ_2 , we first establish that, for any nonnegative real numbers a, b, c and d , $\frac{ax-b}{cx+d}$ is increasing in x . This is because $\left(\frac{ax-b}{cx+d}\right)' = \frac{ad+bc}{(cx+d)^2} \geq 0$. Accordingly, the numerator of ψ_2 can be written as $a\phi - b$ where $a = y_{00} \geq 0$ and $b = \frac{k}{r(a_H - a_L)} > 0$. the denominator of ψ_2 can be written as $c\phi + d$ where $c = y_{11} + y_{00} - 1 \geq 0$ and $d = (1 - y_{11})\beta \geq 0$. Therefore, ψ_2 is increasing in ϕ . □

B.4 Proof of Proposition 2 and Proposition E.1

Proof. We characterize the optimal disclosure policy by allowing both $b_m > b_h$ (or equivalently, $\beta > 1/2$) and $b_m \leq b_h$ (or equivalently, $\beta \leq 1/2$), thereby proving Proposition 2 and Proposition E.1, respectively.

The expected monetary profit as a function v can be written as

$$\Pi(v) = \begin{cases} (r - c)a_L & \text{if } v < \psi_1 \\ A_I v + B_I & \text{if } \psi_1 \leq v \leq \psi_2 \\ A_N v + B_N & \text{if } v > \psi_2 \end{cases} \quad (\text{B.16})$$

where A_I , B_I , A_N and B_N are defined in (B.2), (B.3), (B.5) and (B.6), respectively, and ψ_1 and ψ_2 are given in (B.14)-(B.15). As shown in Kamenica and Gentzkow (2011), under the optimal information design, the maximum expected profit is attained on the upper concave envelope of function $\Pi(v)$, denoted by $\hat{\Pi}(v)$. The upper concave envelope of $\Pi(v)$ is the minimum concave function, which is (weakly) larger than $\Pi(v)$. Formally, the upper concave envelope of $\Pi(v)$ is defined as

$$\hat{\Pi}(v) = \sup\{x : (v, x) \in \text{co}(\Pi)\}$$

where $\text{co}(\Pi)$ is the convex hull of Π .

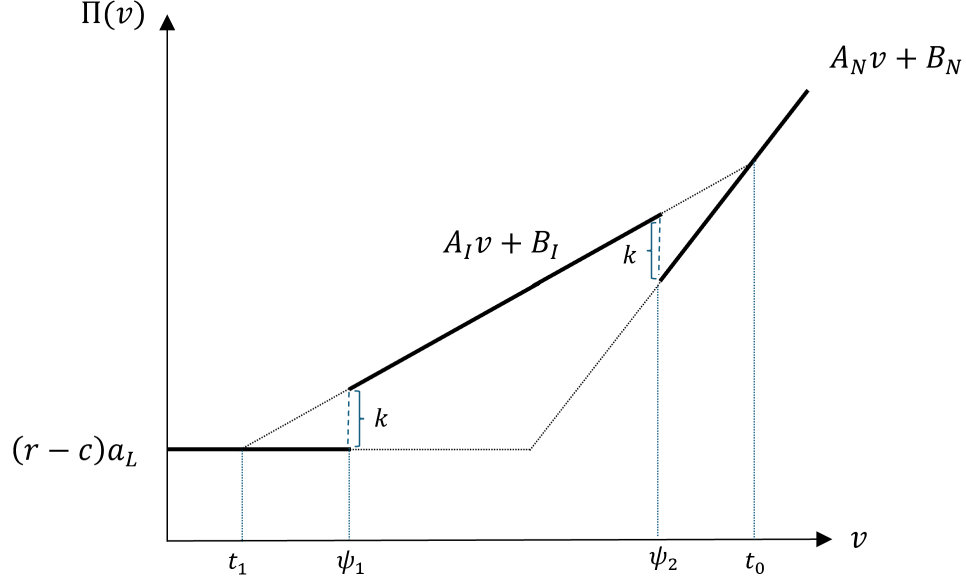


Figure B.1: Illustration of key thresholds used in the analysis

Figure B.1 illustrates $\Pi(v)$ and the useful thresholds $t_1 < \psi_1 \leq \psi_2 < t_0$. Recall that we focus on the range of $\hat{k}$ that excludes the trivial cases (N) and (A) in Lemma 2, which ensures that there exist $\psi_1 \leq \psi_2$, $\psi_1 < 1$ and $\psi_2 > 0$. As shown Figure B.1, $A_I v + B_I$ crosses $(r-c)a_L$ at $v = t_1$ and crosses $A_N v + B_N$ at $v = t_0$, and that $A_I \psi_1 + B_I = (r-c)a_L + k$ and $A_I \psi_2 + B_I = A_N \psi_2 + B_N + k$. Moreover, the following relations will be useful:

$$t_1 \leq 0 \iff \phi \leq 1 - \beta \text{ and } t_0 < 1 \iff \phi < \beta, \quad (\text{B.17})$$

$$\psi_1 \leq \lambda \iff \phi \leq u_n = b_n - \hat{k}/\lambda, \quad (\text{B.18})$$

$$\lambda \leq \psi_2 \iff \phi \geq l_n = \frac{\hat{k} + \lambda(1 - b_n)}{1 - \lambda}. \quad (\text{B.19})$$

In what follows, we examine cases (i), (ii) and (iii) in Lemma 2, respectively. For each case, we characterize the structure of optimal messaging policies based on whether $\beta > 1/2$ (as in Proposition 2) or $\beta \leq 1/2$ (as in Proposition E.1). Taken together, the two propositions will be proved.

Cases (i): $\phi < \min(g_{10}^{-1}(\hat{k}), g_{21}^{-1}(\hat{k}))$ In Case (i), we have $\psi_1 \leq 0$ and $0 < \psi_2 \leq 1$. Then, $\Pi(v)$ reduces to

$$\Pi(v) = \begin{cases} A_I v + B_I & \text{if } v \leq \psi_2, \\ A_N v + B_N & \text{if } v > \psi_2. \end{cases} \quad (\text{B.20})$$

Two cases can emerge. See Figure B.2 for an illustration.

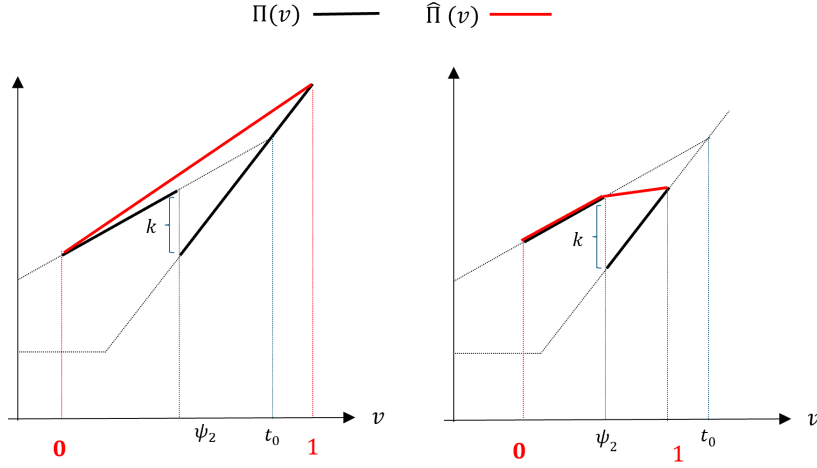


Figure B.2: Graphic illustration of the upper concave envelope for Case (i)

- (i-a)** First, if $t_0 \leq 1$, the upper concave envelope $\hat{\Pi}$ is a linear function connecting $(0, B_I)$ and $(1, A_N + B_N)$. In this case, full revelation of X_m is optimal for any prior λ .
- (i-b)** Second, if $t_0 > 1$, $\hat{\Pi}$ is a piecewise linear function with the first segment connecting $(0, B_I)$ and $(\psi_2, A_I \psi_2 + B_I)$ and the second segment connecting $(\psi_2, A_I \psi_2 + B_I)$ and $(1, A_N + B_N)$. In this case, partial disclosure is optimal. Furthermore,

- If $\lambda > \psi_2$, it is optimal to induce posterior $v = \psi_2$ with probability $\frac{1-\lambda}{1-\psi_2}$ and induce $v = 1$ with probability $\frac{\lambda-\psi_2}{1-\psi_2}$. To achieve such a distribution of v , the information structure is given by

$$\begin{aligned} \eta^*(1|1) &= \frac{\lambda - \psi_2}{1 - \psi_2} \frac{1}{\lambda}, & \eta^*(0|1) &= 1 - \eta^*(1|1); \\ \eta^*(0|0) &= 1, & \eta^*(1|0) &= 0. \end{aligned} \tag{D'}$$

- If $\lambda \leq \psi_2$, $\hat{\Pi}$ coincides with the original, linear function Π , and so multiple solutions exist. It is optimal to induce posterior $v = 0$ with probability $\frac{\lambda}{\psi_2}$ and induce $v = \psi_2$ with probability $1 - \frac{\lambda}{\psi_2}$.¹⁸ This posterior distribution can be achieved by the following

¹⁸In this case, multiple solutions exist to achieve the upper concave envelope. For instance, providing no information about X_m also achieves the upper concave envelope. Throughout the paper, we break the tie by choosing the more informative policy.

information structure.

$$\begin{aligned} \eta^*(1|1) &= 1, \quad \eta^*(0|1) = 0; \\ \eta^*(0|0) &= \left(1 - \frac{\lambda}{\psi_2}\right) \frac{1}{1 - \lambda}, \quad \eta^*(1|0) = 1 - \eta^*(0|0). \end{aligned} \tag{E}$$

Now we show that the above can be simplified according whether $\beta > 1/2$ or $\beta \leq 1/2$.

For Case (i) with $\beta > 1/2$, only Case (i-a) can occur, because $\beta > 1/2$ implies $t_0 \leq 1$. To see this, note that under case (i), $\phi \leq g_{10}^{-1}(\hat{k}) = \frac{(1-\beta)(1-y_{00})-\hat{k}}{1-y_{00}}$, which implies $\phi < 1 - \beta$. Given $\beta > 1/2$, it follows that $\phi < \beta \iff t_0 < 1$ by (B.17). This proves Proposition 2(i).

For Case (i) with $\beta \leq 1/2$, both cases (i-a) and (i-b) can occur. Invoking the fact that $\phi < \beta \iff t_0 < 1$ and $\lambda \leq \psi_2 \iff \phi \geq l_n$ proves Proposition E.1(i).

Case (ii): $\phi \geq \max(g_{10}^{-1}(\hat{k}), g_{21}^{-1}(\hat{k}))$. We have $\psi_1 > 0$ and $1 \leq \psi_2$. Then, $\Pi(v)$ reduces to

$$\Pi(v) = \begin{cases} (r - c)a_L & \text{if } v \leq \psi_1, \\ A_I v + B_I & \text{if } v > \psi_1. \end{cases} \tag{B.21}$$

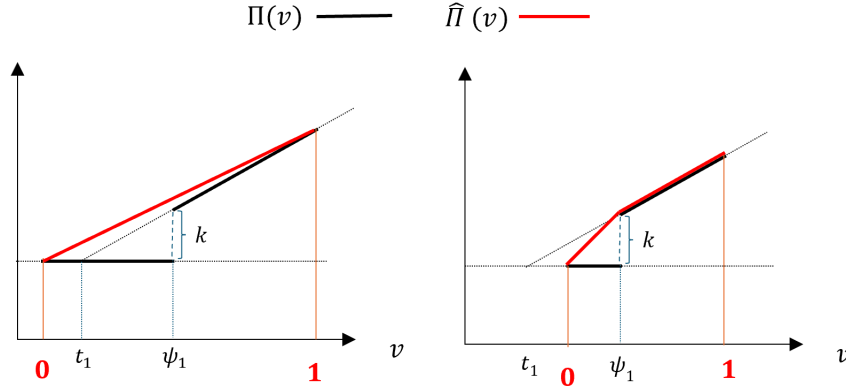


Figure B.3: Graphic illustration of the upper concave envelope for Case (ii)

Two cases can emerge.

(ii-a) First, if $t_1 \geq 0$, then $\hat{\Pi}$ is a linear function crossing points $(0, (r - c)a_L)$ and $(1, A_I + B_I)$. In this case, full revelation of X_m is optimal.

(ii-b) Second, if $t_1 < 0$, then $\hat{\Pi}$ is piece-wise linear function with the first segment connecting $(0, (r - c)a_L)$ and $(\psi_1, (r - c)a_L + k)$ and the second segment connecting $(\psi_1, (r - c)a_L + k)$ and $(1, A_I + B_I)$. In this case, partial revelation is optimal.

- If prior $\lambda \leq \psi_1$, it is optimal to induce posterior $v = 0$ with probability $\frac{\psi_1 - \lambda}{\psi_1}$ and induce $v = \psi_1$ with probability $\frac{\lambda}{\psi_1}$. The information structure resulting in this posterior distribution is as follows.

$$\begin{aligned}\eta^*(1|1) &= 1, & \eta^*(0|1) &= 0 \\ \eta^*(0|0) &= \frac{\psi_1 - \lambda}{\psi_1} \frac{1}{1 - \lambda}, & \eta^*(1|0) &= 1 - \eta^*(0|0).\end{aligned}\tag{E'}$$

- If prior $\lambda > \psi_1$, it is optimal to induce posterior $v = \psi_1$ with probability $\frac{1 - \lambda}{1 - \psi_1}$ and induce $v = 1$ with probability $\frac{\lambda - \psi_1}{1 - \psi_1}$. The information structure resulting in this posterior distribution is given by

$$\begin{aligned}\eta^*(1|1) &= \frac{\lambda - \psi_1}{1 - \psi_1} \frac{1}{\lambda}, & \eta^*(0|1) &= 1 - \eta^*(1|1) \\ \eta^*(0|0) &= 1, & \eta^*(1|0) &= 0.\end{aligned}\tag{D}$$

Now consider the cases of $\beta > 1/2$ and $\beta \leq 1/2$, respectively.

For Case (ii) with $\beta > 1/2$, only case (ii-a) occurs because $\beta > 1/2$ implies $t_1 \geq 0$. To see this, under case (ii), $\phi > g_{21}^{-1}(\hat{k}) = \frac{\hat{k} + \beta(1 - y_{11})}{1 - y_{11}}$ which implies $\phi > \beta$. Given $\beta > 1/2$, it follows that $\phi > 1 - \beta \iff t_1 > 0$ by (B.17). This proves Proposition 2(ii).

For Case (ii) with $\beta \leq 1/2$, both cases (ii-a) and (ii-b) can occur. Using $t_1 < 0 \iff \phi < 1 - \beta$ and $\lambda > \psi_1 \iff \phi < u_n$, we obtain Proposition E.1(ii).

Case (iii): $g_{10}^{-1}(\hat{k}) < \phi < g_{21}^{-1}(\hat{k})$ In this case, $0 < \psi_1 \leq \psi_2 < 1$, and so $\Pi(v)$ consists of three pieces.

$$\Pi(v) = \begin{cases} (r - c)a_L & \text{if } v < \psi_1, \\ A_I v + B_I & \text{if } \psi_1 \leq v \leq \psi_2 \\ A_N v + B_N & \text{if } v > \psi_2 \end{cases}\tag{B.22}$$

- When $0 \leq t_1$, as illustrated in Figure B.4, we have the following cases.

(iii-a) If $\frac{A_I \psi_2 + B_I - (r - c)a_L}{\psi_2} \leq A_N + B_N - (r - c)a_L$, the upper concave envelope $\hat{\Pi}$ is a linear function connecting $(0, (r - c)a_L)$ and $(1, A_N + B_N)$. For this case, full information revelation is optimal.

(iii-b) If $\frac{A_I \psi_2 + B_I - (r - c)a_L}{\psi_2} > A_N + B_N - (r - c)a_L$, the upper concave envelope $\hat{\Pi}$ is a piece-wise linear function connecting $(0, (r - c)a_L)$, $(\psi_2, A_N \psi_2 + B_N + k)$ and then $(1, A_N + B_N)$.

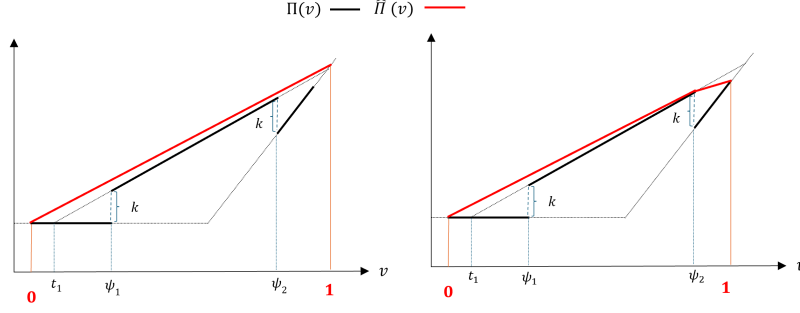


Figure B.4: Graphic illustration of the upper concave envelope for Case (iii) with $t_1 \geq 0$

For the optimal information structure, if $0 < \lambda \leq \psi_2$, it is optimal to induce $v = 0$ with probability $\frac{\psi_2 - \lambda}{\psi_2}$ and induce $v = \psi_2$ with probability $\frac{\lambda}{\psi_2}$. As discussed above, this posterior distribution results from information structure (E). If $\lambda > \psi_2$, it is optimal to induce $v = \psi_2$ with probability $\frac{1 - \lambda}{1 - \psi_2}$ and induce $v = 1$ with probability $\frac{\lambda - \psi_2}{1 - \psi_2}$. This posterior distribution results from (D').

- When $t_1 < 0$, as illustrated in Figure B.5, we have the following cases.

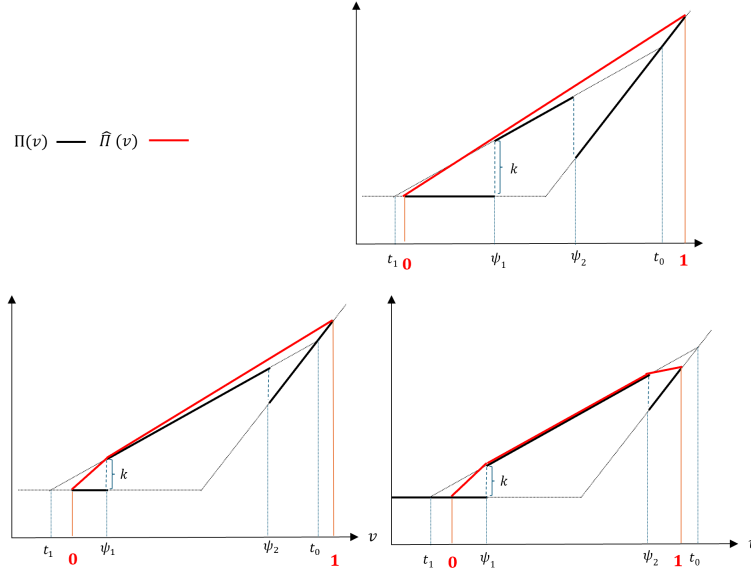


Figure B.5: Graphic illustration of the upper concave envelope for Case (iii) with $t_1 < 0$

- (iii-c) If $\frac{A_I \psi_1 + B_I - (r-c)a_L}{\psi_1} \leq A_N + B_N - (r-c)a_L$ and $t_0 \leq 1$, then the upper concave envelope $\hat{\Pi}$ is a linear function connecting $(0, (r-c)a_L)$ and $(1, A_N + B_N)$. For this case, full information revelation is optimal.
- (iii-d) If $\frac{A_I \psi_1 + B_I - (r-c)a_L}{\psi_1} > A_N + B_N - (r-c)a_L$ and $t_0 \leq 1$, the upper concave envelope $\hat{\Pi}$ is a

piecewise linear function connecting $(0, (r-c)a_L)$, $(\psi_1, (r-c)a_L + k)$ and $(1, A_N + B_N)$.

If $\lambda \leq \psi_1$, then $\hat{\Pi}$ is attained by inducing $v = 0$ with probability $\frac{\psi_1 - \lambda}{\psi_1}$ and $v = \psi_1$ with probability $\frac{\lambda}{\psi_1}$. This posterior distribution results from the information structure (E') as characterized above.

If $\lambda > \psi_1$, then $\hat{\Pi}$ is attained by inducing $v = \psi_1$ with probability $\frac{1-\lambda}{1-\psi_1}$ and $v = 1$ with probability $\frac{\lambda - \psi_1}{1-\psi_1}$. This posterior distribution results from the information structure (D).

(iii-e) If $t_0 > 1$, the upper concave envelope $\hat{\Pi}$ is a piecewise linear function consisting of three pieces, connecting $(0, (r-c)a_L)$, $(\psi_1, (r-c)a_L + k)$, $(\psi_2, A_N\psi_2 + B_N + k)$ and $(1, A_N + B_N)$. This leads to three different cases.

If $\lambda \leq \psi_1$, it is optimal to induce $v = 0$ with probability $\frac{\psi_1 - \lambda}{\psi_1}$ and induce $v = \psi_1$ with probability $\frac{\lambda}{\psi_1}$. The corresponding information structure is given in (E').

If $\psi_1 < \lambda \leq \psi_2$, it is optimal to induce $v = \psi_1$ with probability $\frac{\psi_2 - \lambda}{\psi_2 - \psi_1}$ and induce $v = \psi_2$ with probability $\frac{\lambda - \psi_1}{\psi_2 - \psi_1}$. As such, it is optimal to obfuscate both states. The information structure resulting in such a distribution of v can be constructed as follows:

$$\begin{aligned} \eta^*(1|1) &= \frac{\psi_2}{\lambda} \frac{\lambda - \psi_1}{\psi_2 - \psi_1}, & \eta^*(0|1) &= 1 - \eta^*(1|1) \\ \eta^*(0|0) &= \frac{1 - \psi_1}{1 - \lambda} \frac{\psi_2 - \lambda}{\psi_2 - \psi_1}, & \eta^*(1|0) &= 1 - \eta^*(0|0). \end{aligned} \tag{D\&E}$$

If $\lambda > \psi_2$, it is optimal to induce $v = \psi_2$ with probability $\frac{1-\lambda}{1-\psi_2}$ and induce $v = 1$ with probability $\frac{\lambda - \psi_2}{1-\psi_2}$, which results from the information structure as characterized in (D').

We now reorganize the above cases based on whether $\beta > 1/2$ or $\beta \geq 1/2$.

For Case (iii) with $\beta > 1/2$ The value of ϕ can fall into one of three intervals: $(g_{10}^{-1}(\hat{k}), 1 - \beta)$, $[1 - \beta, \beta]$, and $(\beta, g_{21}^{-1}(\hat{k}))$.

(1) When $\phi \in (g_{10}^{-1}(\hat{k}), 1 - \beta)$, then $t_1 < 0$ and $t_0 < 1$. we have two subcases:

- If $\frac{A_I\psi_1 + B_I - (r-c)a_L}{\psi_1} \leq A_N + B_N - (r-c)a_L$, it leads to Case (iii-c) in which full information is optimal;
- if $\frac{A_I\psi_1 + B_I - (r-c)a_L}{\psi_1} > A_N + B_N - (r-c)a_L$, we have Case (iii-d) in which partial information is optimal with information structure (E') if $\lambda \leq \psi_1$ or (D) if $\lambda > \psi_1$.

We can reduce the condition

$$\frac{A_I \psi_1 + B_I - (r - c)a_L}{\psi_1} \leq A_N + B_N - (r - c)a_L. \quad (\text{K1})$$

as follows. Notice that $\psi_1 = \frac{h_n(\phi)}{h_d(\phi)}$ where $h_n(\phi)$ is a linear increasing function of ϕ and $h_d(\phi)$ is a linear decreasing function of ϕ and both are positive when $\psi_1 > 0$. Moreover, by the definitions of ψ_1 and A_N and B_N , we have

$$A_I \psi_1 + B_I - (r - c)a_L = k \quad (\text{B.23})$$

and

$$A_N + B_N - (r - c)a_L = r(a_H - a_L)(y_{11}(1 - \beta) + \beta - \phi). \quad (\text{B.24})$$

Substituting (B.23)-(B.24) and multiplying both sides of (K1) by $\frac{h_n(\phi)}{r(a_H - a_L)}$, we have

$$(\text{K1}) \text{ holds} \iff \kappa_1(\phi) \leq 0$$

where

$$\kappa_1(\phi) = \hat{k}h_d(\phi) - (y_{11}(1 - \beta) + \beta - \phi)h_n(\phi).$$

Because $h_n(\phi)$ is linearly increasing, $\kappa_1(\phi)$ is a quadratic convex. With the equivalency between $\kappa_1(\phi) \leq 0$ and (K1), we can check whether $\kappa_1(\phi) \leq 0$ by evaluating whether (K1) holds at two limiting points:

- As $\phi \rightarrow g_{10}^{-1}(\hat{k})$, $\psi_1 \rightarrow 0$. Consequently, the left-hand side (LHS) of (K1), which equals $\frac{k}{\psi_1}$, goes to $+\infty$. So, (K1) must be violated and so $\kappa_1(\phi) > 0$, as $\phi \rightarrow g_{10}^{-1}(\hat{k})$.
- As $\phi \rightarrow 1 - \beta$, $t_1 \rightarrow 0$. By the definition of t_1 , at the limiting case $t_1 = 0$, $B_I = (r - c)a_L$, and consequently, the LHS of (K1) reduces to A_I and the right-hand side (RHS) of (K1) equals $A_N + B_N - B_I$. Note that $A_I + B_I < A_N + B_N$ iff $t_0 < 1$ (which is easy to see from Figure B.1). Because $t_0 < 1$ in the current case, (K1) strictly holds and $\kappa_1(\phi) < 0$ as $\phi \rightarrow 1 - \beta$.

Therefore, as ϕ increases from $g_{10}^{-1}(\hat{k})$ to $1 - \beta$, $\kappa_1(\phi)$ changes its sign from positive to negative. By convexity, the sign of $\kappa_1(\phi)$ can change from positive to negative at most once. Thus, there exists a unique threshold $\hat{\phi}_1 \in (g_{10}^{-1}(\hat{k}), 1 - \beta)$ such that $\kappa_1(\phi) > 0$, i.e., (K1) is violated, iff

$$\phi < \hat{\phi}_1.$$

Furthermore, recall that $\psi_1 \leq \lambda \iff \phi \leq u_n$. At $\phi = u_n$, we have

$$\left(\frac{A_I \psi_1 + B_I - (r-c)a_L}{\psi_1} \right) - (A_N + B_N - (r-c)a_L) = (a_H - a_L)(b_h - b_m)$$

which is negative when $\beta > 1/2$. Thus, (K1) holds and $\kappa_1(\phi) < 0$ at $\phi = u_n$. It follows that $\hat{\phi}_1 < u_n$, and therefore $\phi < \hat{\phi}_1$ implies $\phi < u_n$ and hence $\psi_1 < \lambda$.

Consequently, if $g_{10}^{-1}(\hat{k}) < \phi < \hat{\phi}_1$, then $\psi_1 < \lambda$ and (K1) does not hold, and so partial information disclosure with (D) is optimal; if $\hat{\phi}_1 \leq \phi < 1 - \beta$, full information is optimal.

(2) If $\phi \in [1 - \beta, \beta]$, then $t_1 \geq 0$ and $t_0 \leq 1$. In this case, it is easy to see that regardless of k , the upper concave envelope must be a linear line connecting $(0, (r-c)a_L)$ and $(1, A_N + B_N)$, as this linear line is above $A_I v + B_I$ for any v . Hence, full information disclosure is optimal.

(3) If $\phi \in (\beta, g_{21}^{-1}(\hat{k}))$, then $t_0 > 1$ and $t_1 > 0$. Two subcases can emerge.

- If $\frac{A_I \psi_2 + B_I - (r-c)a_L}{\psi_2} \leq A_N + B_N - (r-c)a_L$, we reach Case (iii-a) in which full information is optimal;
- if $\frac{A_I \psi_2 + B_I - (r-c)a_L}{\psi_2} > A_N + B_N - (r-c)a_L$, we reach Case (iii-b) and partial information disclosure is optimal with (E) if $\lambda \leq \psi_2$ or with (D') if $\lambda > \psi_2$.

We analyze whether condition

$$\frac{A_I \psi_2 + B_I - (r-c)a_L}{\psi_2} \leq A_N + B_N - (r-c)a_L \quad (\text{K2})$$

holds as ϕ increases from β to $g_{21}^{-1}(\hat{k})$. By the definitions of ψ_2 and B_N , we have $A_I \psi_2 + B_I = A_N \psi_2 + B_N + k$ and $B_N - (r-c)a_L = (a_H - a_L)(r(1 - y_{00})(1 - \beta) - c)$. Moreover, note that $\psi_2 = \frac{h_n(\phi)}{h_d(\phi)}$ where both $h_n(\phi)$ and $h_d(\phi)$ are linear increasing functions and positive within the relevant parameter range ($0 < \psi_2 < 1$). Substituting $A_I \psi_2 + B_I = A_N \psi_2 + B_N + k$ and $B_N - (r-c)a_L = (a_H - a_L)(r(1 - y_{00})(1 - \beta) - c)$, multiplying both sides by $\frac{h_n(\phi)}{r(a_H - a_L)}$ and rearranging the terms, we have

$$(\text{K2}) \text{ holds} \iff \kappa_2(\phi) \leq 0$$

where

$$\kappa_2(\phi) = (h_d(\phi) - h_n(\phi))((1 - y_{00})(1 - \beta) - \phi) + \hat{k}h_d(\phi).$$

By the definition of ψ_2 , $h_d(\phi) - h_n(\phi)$ is a linear function of ϕ with coefficient $-(1 - y_{11}) \leq 0$. Therefore, $\kappa_2(\phi)$ is a quadratic convex function. With the equivalency of (K2) and $\kappa_2(\phi) \leq 0$, we examine the sign of $\kappa_2(\phi)$ by evaluating whether (K2) holds at two limiting points.

- As $\phi \rightarrow g_{21}^{-1}(\hat{k})$, $\psi_2 \rightarrow 1$, and (K2) reduces to $A_I + B_I \leq A_N + B_N$. However, in Case (3), $t_0 > 1$, which implies that $A_I + B_I > A_N + B_N$. Thus, (K2) is violated and $\kappa_2(\phi) > 0$ as $\phi \rightarrow g_{21}^{-1}(\hat{k})$.
- As $\phi \rightarrow \beta$, $t_0 \rightarrow 1$ such that $A_N + B_N \rightarrow A_I + B_I$. As a result, the RHS of (K2) approaches $A_I + B_I - (r - c)a_L$. On the other hand, the LHS of (K2) equals $A_I + \frac{B_I - (r - c)a_L}{\psi_2}$. Notice that $A_I t_1 + B_I = (r - c)a_L$, and so $B_I - (r - c)a_L = -A_I t_1$. For Case (3), $t_1 > 0$ and so $B_I - (r - c)a_L < 0$. Consequently, since $\psi_2 < 1$, the LHS is smaller than the RHS, that is, (K2) strictly holds and $\kappa_2(\phi) < 0$ as $\phi \rightarrow \beta$.

Therefore, $\kappa_2(\phi) < 0$ as $\phi \rightarrow \beta$ and $\kappa_2(\phi) > 0$ as $\phi \rightarrow g_{21}^{-1}(\hat{k})$. By convexity, $\kappa_2(\phi)$ can change its sign from negative to positive at most once. Thus, there exists a unique threshold $\hat{\phi}_2 \in (\beta, g_{21}^{-1}(\hat{k}))$ such that $\kappa_2(\phi) \leq 0$ and (K2) holds iff $\phi \leq \hat{\phi}_2$. Furthermore, recall that $\lambda \leq \psi_2 \iff \phi \geq l_n$. At $\phi = l_n$, we have

$$\left(\frac{A_I \psi_2 + B_I - (r - c)a_L}{\psi_2} \right) - (A_N + B_N - (r - c)a_L) = (a_H - a_L)(b_h - b_m)$$

which is negative when $\beta > 1/2$. Thus, at $\phi = l_n$, (K2) holds and $\kappa_2(\phi) \leq 0$. It follows that $l_n \leq \hat{\phi}_2$ and so $\phi > \hat{\phi}_2$ implies $\phi > l_n$ and $\lambda < \psi_2$.

Therefore, if $\beta < \phi \leq \hat{\phi}_2$, (K2) holds and full information is optimal; if $\hat{\phi}_2 < \phi < g_{21}^{-1}(\hat{k})$, then (K2) does not hold and $\lambda < \psi_2$, implying that partial information disclosure with (E) is optimal.

Combining the three cases completes the proof of Proposition 2(iii).

For Case (iii) with $\beta \leq 1/2$ For the case of $\beta \leq 1/2$, the value of ϕ can fall into three intervals, $(g_{10}^{-1}(\hat{k}), \beta]$, $(\beta, 1 - \beta)$, and $[1 - \beta, g_{21}^{-1}(\hat{k}))$.

- (1) If $\phi \in (g_{10}^{-1}(\hat{k}), \beta]$, then $t_0 \leq 1$ and $t_1 < 0$. Two subcases could emerge.

- If $\frac{A_I\psi_1+B_I-(r-c)a_L}{\psi_1} \leq A_N + B_N - (r-c)a_L$, i.e., (K1) defined earlier holds, it leads to Case (iii-c) in which full information is optimal;
- if $\frac{A_I\psi_1+B_I-(r-c)a_L}{\psi_1} > A_N + B_N - (r-c)a_L$, i.e., (K1) does not hold, we have Case (iii-d) in which partial information is optimal with (E') if $\lambda \leq \psi_1$ or with (D) if $\lambda > \psi_1$.

In the earlier analysis, we have established (K1) holds iff $\kappa_1(\phi) \leq 0$ where $\kappa_1(\phi)$ is a quadratic convex function. Consider the sign of $\kappa_1(\phi)$ at three points: $\phi = g_{10}^{-1}(\hat{k})$, β and $1 - \beta$. We have shown that $\kappa_1(\phi) > 0$ as $\phi \rightarrow g_{10}^{-1}(\hat{k})$, and $\kappa_1(\phi) < 0$ at $\phi = 1 - \beta$. Now consider $\phi = \beta$, in which case $t_0 = 1$ and so $A_I + B_I = A_N + B_N$. Thus, (K1) reduces to $\frac{B_I-(r-c)a_L}{\psi_1} \leq B_I - (r-c)a_L$. Given $t_1 < 0$ in the current case, we have $B_I - (r-c)a_L > 0$ (as $A_I t_1 + B_I = (r-c)a_L$). For $0 < \psi_1 < 1$, it follows that (K1) does not hold and equivalently $\kappa_1(\phi) > 0$ at $\phi = \beta$. By convexity, $\kappa_1(\phi)$ can change its sign from positive to negative at most once. Given $g_{10}^{-1}(\hat{k}) < \beta \leq 1 - \beta$, it follows that for any $\phi \in (g_{10}^{-1}(\hat{k}), \beta]$, $\kappa_1(\phi) > 0$, i.e., (K1) is violated. So, partial information disclosure is optimal as described in the second subcase, whereas the first subcase cannot emerge. Unlike the case of $\beta > 1/2$, here we cannot rule out the possibility of $\lambda \leq \psi_1$ or $\lambda > \psi_1$. By $\psi_1 \leq \lambda \iff \phi \leq u_n$, (D) is optimal if $\phi < u_n$, and (E') is optimal if $\phi \geq u_n$.

- (2) If $\beta < \phi < 1 - \beta$, then $t_1 < 0$ and $t_0 > 1$. We have Case (iii-e) where the optimal information structure is (E') if $\lambda \leq \psi_1 \iff \phi \geq u_n$, (D&E) if $\psi_1 < \lambda \leq \psi_2 \iff l_n \leq \phi < u_n$, or (D') if $\lambda > \psi_2 \iff \phi < l_n$.
- (3) If $\beta \leq 1 - \beta < \phi$, then $t_0 > 1$ and $t_1 > 0$, and we have two possible subcases:

- if $\frac{A_I\psi_2+B_I-(r-c)a_L}{\psi_2} \leq A_N + B_N - (r-c)a_L$, i.e., (K2) defined earlier holds, we have Case (iii-a) in which full information is optimal;
- if $\frac{A_I\psi_2+B_I-(r-c)a_L}{\psi_2} > A_N + B_N - (r-c)a_L$, i.e., (K2) does not hold, we reach Case (iii-b) where partial information is optimal with (E) if $\lambda \leq \psi_2 \iff \phi \geq l_n$ or with (D') if $\lambda > \psi_2 \iff \phi < l_n$.

The preceding analysis has shown that (K2) holds iff $\kappa_2(\phi) \leq 0$ where κ_2 is a quadratic convex function. Consider the sign of $\kappa_2(\phi)$ at three points: $\phi = \beta, 1 - \beta$ and $g_{21}^{-1}(\hat{k})$. We have shown that $\kappa_2(\phi) < 0$ at $\phi = \beta$ and $\kappa_2(\phi) > 0$ as $\phi \rightarrow g_{21}^{-1}(\hat{k})$. Now consider $\phi = 1 - \beta$, in which $t_1 = 0$ and so $B_I = (r-c)a_L$. Consequently, (K2) can be reduced to $A_I \leq A_N + B_N - B_I$. Because $A_I + B_I > A_N + B_N$ when $t_0 > 1$, (K2) is violated and equivalently $\kappa_2(\phi) > 0$ at

$\phi = 1 - \beta$. By convexity, $\kappa_2(\phi)$ can change its sign from negative to positive at most once as ϕ increases. Given $\beta \leq 1 - \beta < g_{21}^{-1}(\hat{k})$ and that $\kappa_2(\phi)$ is negative, positive and positive at $\phi = \beta$, $1 - \beta$ and $g_{21}^{-1}(\hat{k})$, respectively, we can conclude that for any $\phi \in [1 - \beta, g_{21}^{-1}(\hat{k}))$, $\kappa_2(\phi) > 0$ and thus (K2) does not hold. Therefore, partial information is optimal as described in the second subcase, whereas the first subcase never emerges.

Combining the optimality conditions for each messaging policy completes the proof of Proposition E.1(iii). $\square$

B.5 Proof of Proposition 3

Proof. First, note that in the parameter region of Proposition 2, the threshold u_0 is decreasing in b_m , whereas the threshold l_1 is increasing in b_m .

(i) Given some fixed $\phi \leq u_0$, full disclosure is optimal with $\eta^*(1|1) = 1$. As b_m increases, u_0 becomes smaller, and the optimal policy changes to (D) when $\phi > u_0$.

Under policy (D), $\eta^*(1|1) = \frac{\lambda - \psi_1}{1 - \psi_1} \frac{1}{\lambda}$, which is decreasing in ψ_1 . Differentiating ψ_1 with respect to b_m , we observe that $\frac{\partial \psi_1}{\partial b_m} \geq 0$ iff $\hat{k} - \lambda(b_h - \phi) \leq 0$. On the other hand, (D) is optimal only when $\psi_1 \leq \lambda \iff \phi \leq b_h - \hat{k}/\lambda$. It follows that ψ_1 increases with b_m within the region where (D) is optimal. Thus, $\eta^*(1|1)$ is decreasing in b_m under policy (D).

As established in the proof of Proposition 2, the threshold $\hat{\phi}_1 < 1 - \beta$, whereas $1 - \beta$ is decreasing in b_m and approaches zero when b_m becomes sufficiently close to one. This observation together with Proposition 2(iii) implies that, as b_m continues to increase, the optimal policy will shift from (D) to (F), at which point $\eta^*(1|1)$ jumps to one.

(ii) Given some fixed $\phi > l_1$, l_1 becomes larger as b_m increases. The optimal policy will change from (F) to (E) as l_1 exceeds ϕ .

Under Policy (E), $\eta^*(0|0) = \left(1 - \frac{\lambda}{\psi_2}\right) \frac{1}{1 - \lambda}$, which is increasing in ψ_2 . Observe that $\frac{\partial \psi_2}{\partial b_m} \leq 0$ iff $\hat{k} - (\phi(1 - \lambda) - \lambda(1 - b_h)) \leq 0$. On the other hand, (E) is optimal only when $\lambda \leq \psi_2 \iff \phi \geq \frac{\hat{k} + \lambda(1 - b_h)}{1 - \lambda}$. It follows that ψ_2 is decreasing in b_m , and so is $\eta^*(0|0)$ under policy (E).

As established in the proof of Proposition 2, the threshold $\hat{\phi}_2 > \beta$, whereas β is increasing in b_m and approaches to one when b_m is sufficiently close to one. This observation together with Proposition 2(iii) implies that, as b_m continues to increase, the optimal policy will shift from (E) to (F), at which point $\eta^*(0|0)$ jumps to one.

$\square$

B.6 Proof of Proposition 4

See Appendix D.

B.7 Proof of Proposition 5

See Appendix F.

C Value of Optimal Disclosure Policies

In this appendix, we present numerical results demonstrating that partial disclosure of machine information can yield greater profit improvement than deploying a full-disclosure machine.

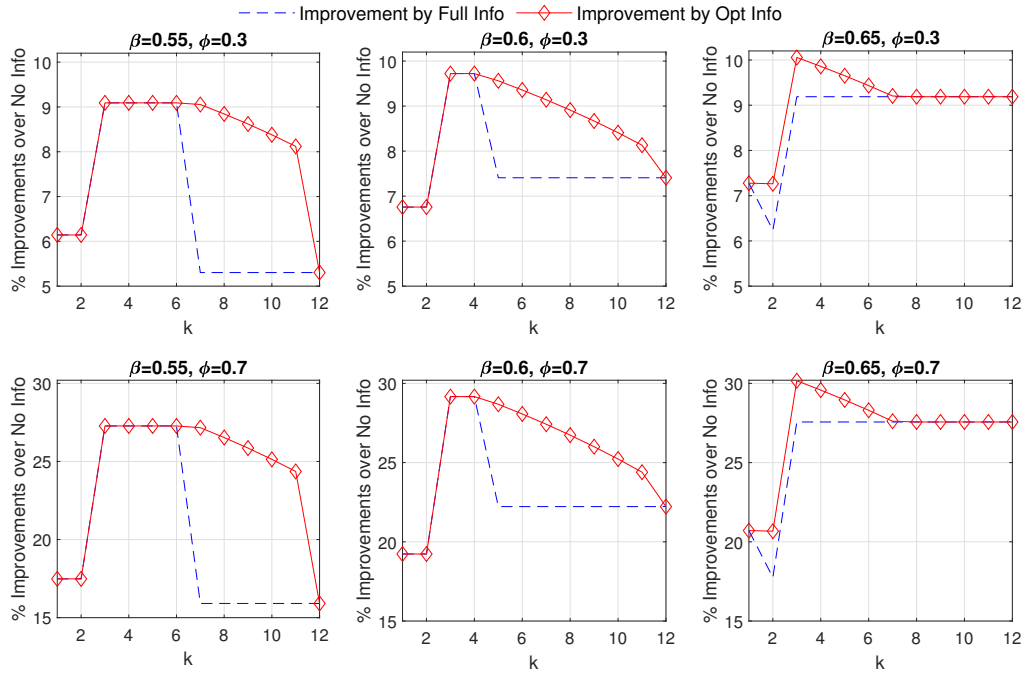


Figure C.1: Profit improvements relative to no machine [Parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$, $b_h = 0.75$ with b_m and c varied according to β and ϕ given in the graph]

Figure C.1 presents a numerical study comparing the improvements in profit, relative to the no-information case, achieved by a full-information disclosure machine and an optimally designed partial-disclosure machine. The gap between the “Full Info” and “Opt Info” curves indicates the additional value partial disclosure brings compared to full disclosure.¹⁹ When partial disclosure

¹⁹The improvements by a full-information machine over the no-machine case is non-monotonic in k for the following reason. When k is small enough, the manager exerts effort both without a machine and under some message realization with a machine. When k is large enough, the manager exerts no effort, whether there is a machine or not. However,

is optimal, it leads to a significant improvement. For example, when $\beta = 0.55$, $\phi = 0.3$ and $k \in [7, 11]$, implementing a full-information machine improves the expected monetary profit by approximately 5.2%. By partially revealing information, the optimally designed machine can boost the improvements to 8 – 9%. The improvements are more significant for $\phi = 0.7$ because for a low-margin product, information has a greater impact on whether one should order a_H (target the high market state) or not. The improvement of partial over full information becomes less significant as β increases, that is, as the machine forecast becomes more accurate.

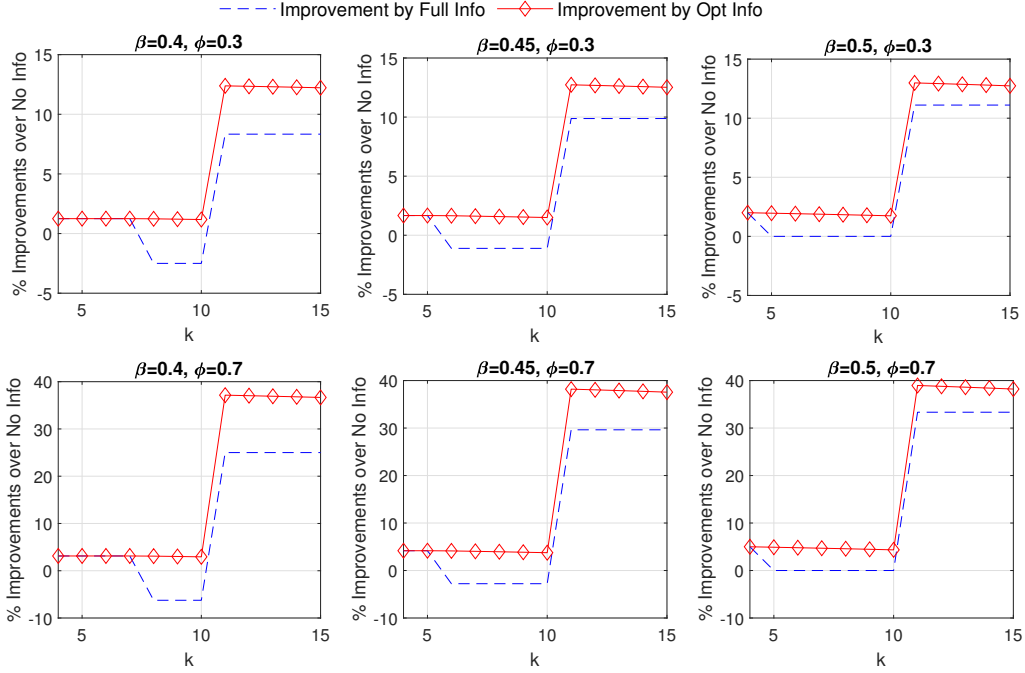


Figure C.2: Profit improvements relative to no machine [Parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$, $b_h = 0.9$ with b_m and c varied according to β and ϕ given in the graph]

Figure C.2 presents another set of results when $b_m \leq b_h$. Again we see that partial disclosure policies, when optimal, can significantly enhance the improvement as compared to full disclosure. When the machine is less accurate than the human ($b_m \leq b_h$), there can be instances (e.g., when $\beta = 0.4$ and $k \in [8, 10]$) where providing full information is worse than the no-machine case because full disclosure dampens human effort. Yet, optimally garbling the machine prediction can overcome this effort-dampening pitfall, thereby enabling the firm to complement human prediction effort and when k is in the middle range between these two cases, the optimal effort choices are indeterminate. The manager may exert effort under some machine message but chooses no effort without a machine, thus leading to larger improvements (as seen in the upward jump for $\beta = 0.55$ and 0.6), or may never exert effort with a machine but exerts effort without one, leading to smaller improvements (as seen in the downward jump at $k = 2$ for $\beta = 0.65$).

achieve a positive improvement.

D Optimized Compensation

In the base model, we assume that the firm follows a performance-based compensation scheme $t(\pi) = w_B + w_P\pi$ where w_B denotes the base salary, and w_P is a commission rate offered to the manager. Given any fixed compensation parameters $\{w_B, w_P\}$, the manager chooses an effort level and order quantity to maximize $w_B + w_P\mathbb{E}[\pi] - \kappa e$, where we use κ to represent the manager's original effort cost before being scaled relative to w_P . Equivalently, the manager maximizes $\mathbb{E}[\pi] - \frac{\kappa}{w_P}e$. Thus, all our results follow by defining a normalized effort cost $k = \kappa/w_P$. Our main findings indicate that, without altering the existing compensation scheme, the firm may be better off obfuscating machine predictions so as to incentivize managerial effort. Alternatively, the firm may optimize the compensation parameters $\{w_B, w_P\}$ to address the incentive issue caused by the machine prediction while providing a full-information machine. Which approach would be better for the firm?

To set a basis for comparison, we consider a status quo in which no prediction machine is available. The firm chooses the optimal compensation parameters, denoted by $\{w_B^N, w_P^N\}$, to incentivize the manager to exert effort (with no machine information). We regard $\{w_B^N, w_P^N\}$ as the *existing* contract that is optimized for the no-machine setting. Then, assuming that a prediction machine becomes available, the firm may (1) retain the existing contract while optimally disclosing the machine information as characterized in our base model, or (2) deploy a full-information machine but re-optimize the compensation parameters with consideration of the manager's effort choice in response to machine predictions. We refer to the former strategy as the information design approach and the latter as the optimal contracting approach. In this appendix, we aim to compare the firm's expected payoffs under these two different approaches.

D.1 Optimal Compensation with No Machine

We first derive the contract parameters that are optimal for the no-machine case. Consider the same setting as in our base model. Given any posterior belief μ regarding the probability of $\theta = H$, the manager effectively chooses between two order sizes a_L and a_H . Consequently, the realized profit

takes one of the three outcomes:

$$\pi = \begin{cases} \pi_O := ra_L - ca_H & \text{if } Q = a_H, \theta = L \\ \pi_L := (r - c)a_L & \text{if } Q = a_L, \theta = L \text{ or } H \\ \pi_H := (r - c)a_H & \text{if } Q = a_H, \theta = H \end{cases} \quad (\text{D.1})$$

where π_O , π_L and π_H represent the realized profits when overstocking, stock and sell a_L units, and stock and sell a_H units, respectively. In line of many retail settings, we assume that actual demand can be censored, that is, when the manager believes the demand likely to be low and orders a_L units, the firm will only observe a_L units of sales without ascertaining whether the actual demand is low or high.²⁰ Given any posterior belief μ , under any compensation scheme, the manager orders a_L units iff $\pi_L \geq \pi_O + \mu(\pi_H - \pi_O)$, or equivalently, $\mu \leq \phi$. Thus, the resulting profit can be written as

$$\max\{\pi_L, \pi_O + \mu(\pi_H - \pi_O)\}.$$

Moreover, the manager's ordering decision is consistent with the profit-maximizing quantity as in Lemma 1. Therefore, it suffices to consider the incentive compatibility (IC) constraint that induces effort.

Recall that $\Pr(X_h = 1) = \lambda$ and, without machine information, the posterior μ conditional on $s_h = 1$ and $s_h = 0$ is given by b_h and $\frac{\lambda(1-b_h)}{1-\lambda}$, respectively. With these probability expressions, any effort-inducing contract must satisfy the following incentive-compatible (IC) constraint:

$$\begin{aligned} & \lambda(w_B + w_P \max\{\pi_L, \pi_O + b_h(\pi_H - \pi_O)\}) + (1 - \lambda)(w_B + w_P \max\{\pi_L, \pi_O + \frac{\lambda(1-b_h)}{1-\lambda}(\pi_H - \pi_O)\}) - \kappa \\ & \geq w_B + w_P \max\{\pi_L, \pi_O + \lambda(\pi_H - \pi_O)\} \end{aligned} \quad (\text{D.2})$$

Note that inducing effort brings value only if the human's prediction is accurate enough to influence the order quantity decision. Specifically, we focus on the parameter range that satisfies the assumption below.

Assumption 3. (Inducing Human Effort Is Valuable) $\frac{\lambda(1-b_h)}{1-\lambda} < \phi < b_h$.

Instead, if $\phi \leq \frac{\lambda(1-b_h)}{1-\lambda}$ or $\phi \geq b_h$, the optimal order quantity will be independent of X_h , and thus there is no value in obtaining a human prediction.

²⁰Consequently, incentive payments contingent on forecast errors may not be implementable since the high demand state is not necessarily verifiable. This further justifies our focus on compensation schemes contingent on realized profits.

In practice, firms are typically required to ensure a minimum salary, regardless of realized profits. The minimum salary requirement can be modeled as a limited liability (LL) constraint such that the compensation payment $t(\pi) \geq \underline{t}$ for all realization of π , where $\underline{t} > 0$ represents the minimum salary level determined by labor regulations or commonly accepted market standards. As such, the firm's choice of $\{w_B, w_P\}$ must satisfy a LL constraint written as

$$w_B + w_P \pi_O \geq \underline{t}. \quad (\text{D.3})$$

We normalize the expected value of the manager's outside option to zero, and so the manager's participation is guaranteed by (D.2) and (D.3).

By choosing w_B and w_P subject to (D.2) and (D.3), the firm maximize the expected payoff

$$(1 - w_P)\mathbb{E}[\pi] - w_B \quad (\text{D.4})$$

where $\mathbb{E}[\pi]$ is equal to the expected profit when the manager exerts effort to learn X_h and chooses the profit-maximizing order quantity accordingly.

Proposition D.1. *With no prediction machine, the optimal compensation $\{w_B^N, w_P^N\}$ that induces managerial effort is characterized below.*

(i) *If $\phi < \lambda$, then $w_P^N = \frac{\kappa}{r(a_H - a_L)((1-\lambda)\phi - \lambda(1-b_h))}$ and $w_B^N = \underline{t} - w_P^N \pi_O$, and the manager obtains an expected payoff $U^N = \underline{t} + w_P^N \lambda(\pi_H - \pi_O)$.*

(ii) *If $\phi \geq \lambda$, then $w_P^N = \frac{\kappa}{\lambda r(a_H - a_L)(b_h - \phi)}$ and $w_B^N = \underline{t} - w_P^N \pi_O$, and the manager obtains an expected payoff $U^N = \underline{t} + w_P^N (\pi_L - \pi_O)$.*

Moreover, there exists an interval $(\underline{\phi}, \bar{\phi})$ such that, for all $\phi \in (\underline{\phi}, \bar{\phi})$, if the firm retains the optimal compensation scheme $\{w_B^N, w_P^N\}$ and deploys a full-information machine, the manager will not exert effort, regardless of the realized machine signals, where $\underline{\phi} = \frac{\lambda(2-b_m-b_h+b_h\lambda-\lambda)}{1-b_h\lambda-b_m\lambda+\lambda^2}$ and $\bar{\phi} = \frac{1-b_h+b_h\lambda}{2-b_h-b_m+\lambda}$.

In both cases, w_P^N is increasing in the manager's effort cost κ , indicating that a higher monetary reward is required as effort becomes more costly. Meanwhile, the base salary term w_B^N is set to satisfy the minimum salary requirement. As is standard in moral hazard settings with a LL constraint, the manager (agent) obtains more rent whereas the firm (principal) incurs a greater agency cost, as the effort cost κ increases.

Notably, we establish that if the firm deploys a full-information machine while retaining the contract (w_B^N, w_P^N) that was optimized prior to the machine introduction, there always exist a moderate range of ϕ (or profit margins) for which the manager will not exert effort. This result indicate that, although our base model analysis is conducted for given contract terms, the core incentive issue persists even when the contract had been optimally designed before the machine was deployed.

Proof of Proposition D.1. It is easy to see that w_B cancels out in the IC constraint. Thus, in the optimal solution, the firm will always set $w_B^N = \underline{t} - w_P\pi$ such that the LL constraint is binding. Consequently, the firm's payoff reduces to $\mathbb{E}[\pi] - \underline{t} - w_P(\mathbb{E}[\pi] - \pi_O)$. Clearly, the optimal w_P^N is given by the smallest w_P that satisfies the IC constraint, which is determined according to each of the two cases below.

- (i) For $\phi < \lambda$, the IC constraint reduces to $\frac{\kappa}{w_P} \leq (1 - \lambda)(\pi_L - \pi_O) - \lambda(1 - b_h)(\pi_H - \pi_O)$. By the expressions of each profit outcomes, the smallest w_P is as given in part (i) of the proposition. Because under the optimal compensation the IC constraint is binding, the manager's expected payoff is given by $U^N = \underline{t} + w_P^N \lambda(\pi_H - \pi_O)$
- (ii) For $\phi \geq \lambda$, the IC constraint reduces to $\frac{\kappa}{w_P} \leq \lambda[b_h(\pi_H - \pi_O) + \pi_O - \pi_L]$. By the expressions of each profit outcomes, the smallest w_P is as given in part (ii) of the proposition. Since under the optimal compensation the IC constraint is binding, the manager's expected payoff is given by $U^N = \underline{t} + w_P^N(\pi_L - \pi_O)$.

Note that, for any given $w_P > 0$, by setting $k = \frac{\kappa}{w_P}$ in the base model, according to Proposition 1 and Figure 1, one can determine if the manager exerts effort with a full-information machine. Graphically, the optimal compensation is such that $k = \kappa/w_P^N$ lies on the boundary of the dashed triangular area in Figure 1, thus incentivizing managerial effort with the smallest w_P . Algebraically, at $k = \kappa/w_P^N$, we always have $\phi = l_n$ for $\phi < \lambda$, and $\phi = u_n$ for $\phi \geq \lambda$. By Proposition 1, the manager never exerts effort with a full-information machine if $u_0 < \phi < l_1$. These observations together imply that if ϕ is between $\underline{\phi}$ and $\bar{\phi}$ where $\underline{\phi}$ is determined by the intersection of thresholds u_0 and l_n , and $\bar{\phi}$ is determined by the intersection of u_n and l_1 . Specifically, $\underline{\phi} = \frac{\lambda(2-b_m-b_h+b_h\lambda-\lambda)}{1-b_h\lambda-b_m\lambda+\lambda^2}$ and $\bar{\phi} = \frac{1-b_h+b_h\lambda}{2-b_h-b_m+\lambda}$. $\square$

D.2 Optimal Compensation with a Full-Information Machine

Suppose now that the firm deploys a full-information machine that fully discloses X_m , and that it can costlessly adjust the existing compensation terms to its advantage.

As in the previous subsection, for any posterior μ , the manager's ordering decision is consistent with the profit-maximizing quantity as in Lemma 1. Therefore, it suffices to consider the incentive compatibility (IC) constraint that induces effort. Since we focus on full disclosure, let us write the posterior as $\mu(X_m, X_h)$ if managerial effort is exerted and as $\mu(X_m, \emptyset)$ otherwise.²¹ Define $U(X_m, e)$ as the manager's payoff given X_m and e . Under any compensation scheme, we have

$$U(X_m, 0) = w_B + w_P \max\{\pi_L, \pi_O + \mu(X_m, \emptyset)(\pi_H - \pi_O)\}$$

and

$$U(X_m, 1) = w_B + w_P \mathbb{E}_{X_h|X_m}[\max\{\pi_L, \pi_O + \mu(X_m, X_h)(\pi_H - \pi_O)\}] - k.$$

To induce effort $e^*(X_m) = 1$, for a given realization of X_m , the IC constraint is written as

$$U(X_m, 1) \geq U(X_m, 0) \tag{D.5}$$

Similar to the information design approach, depending on the value of $\phi = c/r$, the firm may induce $e^*(X_m) = 1$ by imposing the above IC constraint for one of the realizations of X_m .

Additionally, like the no-machine benchmark, to meet the minimum wage requirement, the LL constraint $w_B + w_P \pi_O \geq \underline{t}$ is imposed. By choosing w_B and w_P subject to the IC and LL constraints, the firm maximize its expected payoff

$$(1 - w_P) \mathbb{E}[\pi] - w_B \tag{D.6}$$

where $\mathbb{E}[\pi]$ is the expected newsvendor profit when the manager acquires signal s_h and chooses the optimal order quantity optimally based on each signal realization.

We focus on the case of $\beta > \frac{1}{2}$, and examine two relevant parameter regimes under which the partial disclosure policies (D) and (E) are optimal under the information design approach. For

²¹To recap the expressions of different posteriors, given $X_m = 1$, $\mu(1, \emptyset) = \Pr(\theta = H|X_m = 1) = b_m$, and given $X_m = 0$, $\mu(0, \emptyset) = \Pr(\theta_H|X_m = 0) = \frac{\lambda(1-b_m)}{1-\lambda}$. Conditional on $X_m = 1$, $X_h = 1$ with probability $z(1) = b_m + b_h - 1$; Conditional on $X_m = 0$, $X_h = 1$ with probability $z(0) = \frac{\lambda}{1-\lambda}(2 - b_m - b_h)$. Given any realization of (X_m, X_h) , the posterior belief in $\theta = H$ is given by $\mu(1, 1) = 1$, $\mu(1, 0) = \frac{1-b_h}{2-b_m-b_h}$, $\mu(0, 1) = \frac{1-b_m}{2-b_m-b_h}$, and $\mu(0, 0) = 0$.

simplicity, we refer to these regimes as regime (D) and (E), respectively. As presented in Proposition D.2, the optimal compensation parameters with a full-information machine, denoted by $\{w_B^F, w_P^F\}$, are derived in closed form for each regime.

Proposition D.2. (i) *In regime (D), the optimal compensation scheme is given by*

$$w_P^F = \frac{\kappa(1 - \lambda)}{\lambda r(a_H - a_L)((1 - b_m) - (2 - b_m - b_h)\phi)} \quad (\text{D.7})$$

and $w_B^F = \underline{t} - w_P^F(ra_L - ca_H)$.

(ii) *In regime (E), the optimal compensation scheme is given by*

$$w_P^F = \frac{\kappa}{r(a_H - a_L)((2 - b_m - b_h)\phi - (1 - b_h))} \quad (\text{D.8})$$

and $w_B^F = \underline{t} - w_P^F(ra_L - ca_H)$.

(iii) *Moreover, in both regimes, $w_P^F > w_P^N$.*

A key observation is that, to induce effort with a full-information machine, the firm must raise the commission rate above the existing term w_P^N that was designed for the no-machine setting. This necessitates a higher agency cost under the optimal contracting approach.

Proof of Proposition D.2. (i) When (D) Is Optimal $\implies \frac{c}{r} < 1 - \beta$

For any realized (X_m, X_h) , we have $\mu(1, 1) = 1$, $\mu(1, 0) = \beta$, $\mu(0, 1) = 1 - \beta$, and $\mu(0, 0) = 0$. Accordingly, the optimal order quantities are $Q^*(1, 1) = Q^*(1, 0) = Q^*(0, 1) = a_H$ and $Q^*(0, 0) = a_L$. As under regime (D), it is desirable to induce effort $e^*(X_m) = 1$ iff $X_m = 0$. (There is no value to acquiring X_h given $X_m = 1$ because $Q^*(1, X_h) = a_H$ for all X_h .)

Under these desired actions of the manager, the expected profit is given by

$$\pi_{contract}^D = (p_{11} + p_{10})(\pi_O + \mu(1, \emptyset)(\pi_H - \pi_O)) + p_{01}(\pi_O + \mu(0, 1)(\pi_H - \pi_O)) + p_{00}\pi_L.$$

The firm chooses $w_B, w_P \geq 0$ to maximize the following objective:

$$V_{contract}^D = (1 - w_P)\pi_{contract}^D - w_B. \quad (\text{D.9})$$

To induce effort $e^*(X_m) = 1$ for $X_m = 0$, the IC condition (D.5) reduces to

$$z(0)(\pi_O + (1 - \beta)(\pi_H - \pi_O)) + (1 - z(0))\pi_L - \frac{\kappa}{w_P} \geq \max\{\pi_L, \pi_O + \mu(0, \emptyset)(\pi_H - \pi_O)\}. \quad (\text{D.10})$$

Proposition 2 implies that $\mu(0, \emptyset) \leq \phi$ when policy (D) is optimal, and so $\max \{\pi_L, \pi_O + \mu(0, \emptyset)(\pi_H - \pi_O)\} = \pi_L$.

Clearly, it is optimal for the firm to reduce w_P and w_B such that (D.10) and the LL constraint, $w_B + w_P\pi_O \geq \underline{t}$, are binding, resulting to the optimal compensation

$$w_P^F = \frac{\kappa(1 - \lambda)}{\lambda(a_H - a_L)((1 - b_m)r - (2 - b_m - b_h)c)} \quad (\text{D.11})$$

and

$$w_B^F = \underline{t} - w_P^F(ra_L - ca_H). \quad (\text{D.12})$$

(ii) When (E) Is Optimal $\implies \frac{c}{r} > \beta$

In this case, the optimal order quantities are $Q^*(1, 1) = a_H$ and $Q^*(1, 0) = Q^*(0, 1) = Q^*(0, 0) = a_L$. It is desirable to induce effort $e^*(X_m) = 1$ iff $X_m = 1$. (There is no value to inducing effort given $X_m = 0$, because $Q^*(0, X_h) = a_L$ for all X_h .) Under these desired actions of the manager, the expected profit is equal to

$$\pi_{contract}^E = p_{11}\pi_H + (p_{10} + p_{01} + p_{00})\pi_L.$$

The IC constraint (D.5) reduces to

$$z(1)\pi_H + (1 - z(1))\pi_L - \frac{\kappa}{w_P} \geq \max \{\pi_L, \pi_O + \mu(1, \emptyset)(\pi_H - \pi_O)\} \quad (\text{D.13})$$

Proposition 2 implies that $\mu(1, \emptyset) \geq \phi$ when policy (E) is optimal, and so $\max \{\pi_L, \pi_O + \mu(1, \emptyset)(\pi_H - \pi_O)\} = \pi_O + \mu(1, \emptyset)(\pi_H - \pi_O)$. The optimal compensation is thus given by setting w_P and w_B such that the above IC constraint and the LL constraint are binding:

$$w_P^F = \frac{\kappa}{(a_H - a_L)((2 - b_m - b_h)c - (1 - b_h)r)} \quad (\text{D.14})$$

and

$$w_B^F = \underline{t} - w_P^F(ra_L - ca_H). \quad (\text{D.15})$$

(iii) Comparing w_P^F and w_P^N . Notice that, by Proposition D.1, w_P^N is decreasing in $\phi < \lambda$ and increasing in $\phi \geq \lambda$ (and continuous at $\phi = \lambda$).

In regime (D), w_P^F is increasing in ϕ , and within regime (D), we necessarily have $\underline{\phi} < \phi$ where

$\underline{\phi}$ is given in Proposition D.1. Consequently, for $\phi < \lambda$, $w_P^F - w_P^N$ increases with ϕ . Note that at $\phi = \underline{\phi}$, $w_P^F = w_P^N$. Thus, for all $\phi \in (\underline{\phi}, \lambda)$, $w_P^F > w_P^N$. For $\phi \geq \lambda$, it is straightforward to verify that $\frac{\kappa}{r(a_H - a_L)w_P^F} - \frac{\kappa}{r(a_H - a_L)w_P^N} = -\frac{\lambda}{1-\lambda} \left(\lambda(1 - b_h) + (b_m + b_h - 1 - \lambda)(1 - \phi) \right) < 0$ and hence $w_P^F > w_P^N$.

In regime (E), w_P^F is decreasing in ϕ , and we necessarily have $\phi < \bar{\phi}$ where $\bar{\phi}$ is given in Proposition D.1. For $\phi \in [\lambda, \bar{\phi})$, $w_P^F - w_P^N$ decreases with ϕ and $w_P^F = w_P^N$ at $\phi = \bar{\phi}$; hence, $w_P^F > w_P^N$. For $\phi < \lambda$, it is easy to verify that $\frac{\kappa}{r(a_H - a_L)w_P^F} - \frac{\kappa}{r(a_H - a_L)w_P^N} = -(1-\lambda)(1-b_h) - (b_m + b_h - 1 - \lambda)\phi < 0$ and so $w_P^F > w_P^N$.

□

D.3 Information Design vs Optimal Contracting

Let π_{info} and $\pi_{contract}$ represent the expected total profit under the information design and optimal contracting approaches, respectively. Denote by V_{info} and $V_{contract}$ the corresponding payoffs of the firm. It follows that

$$V_{contract} = (1 - w_P^F)(\pi_{contract} - \pi_O) + \pi_O - \underline{t}$$

and

$$V_{info} = (1 - w_P^N)(\pi_{info} - \pi_O) + \pi_O - \underline{t}.$$

On the one hand, the optimal contracting approach yields higher expected monetary profit, i.e., $\pi_{contract} > \pi_{info}$, because the manager exerts desired effort under full machine information. On the other hand, as shown earlier, $w_P^F > w_P^N$, that is, the optimal contracting approach renders a higher commission rate. As a consequence, the comparison between $V_{contract}$ and V_{info} is nontrivial. Assuming equal priors, we can analytically establish that the information design approach outperforms the optimal contract approach when the effort cost κ is large. As such, when augmenting the manager with a prediction machine, the firm may be better optimally obfuscating machine information without altering the existing compensation terms, than fully disclosing machine information while re-optimizing the compensation terms.

Proposition D.3. *Assume the prior $\lambda = \frac{1}{2}$. $V_{info} > V_{contract}$ if κ exceeds a threshold $\hat{\kappa}$.*

Proof of Proposition D.3. In the proof, we use subscript “contract” and “info” and superscript $i \in \{D, E\}$ to indicate the total profit and the firm’s payoff under each approach and under each policy regime i .

Regime D. First, we consider the regime where policy (D) is optimal for the information design

approach. The existing (no-machine) contract parameter w_P^N is given by case (i) of Proposition D.1. From the concavification approach of our analysis, the expected profit under information disclosure policy (D) can be written as

$$\pi_{info}^D = \frac{1-\lambda}{1-\psi_1}(A_I\psi_1 + B_I) + \frac{\lambda-\psi_1}{1-\psi_1}(A_N + B_N)$$

where A_I, B_I, A_N and B_N are as defined in (B.2), (B.3), (B.5) and (B.6). Substituting w_P^N and π_{info}^D yields the expression of V_{info}^D .

Given full machine information and the optimal contract to induce effort $e^*(0) = 1$, the resulting expected monetary profit is given by

$$\pi_{contract}^D = (p_{11} + p_{10})(\pi_O + \mu(1, \emptyset)(\pi_H - \pi_O)) + p_{01}(\pi_O + \mu(0, 1)(\pi_H - \pi_O)) + p_{00}\pi_L.$$

Substituting $\pi_{contract}^D$ and w_P^D derived in Proposition D.2 yields the expression of $V_{contract}^D$.

Observe that $V_{info}^D - V_{contract}^D$ is a linear function of κ . Then, it suffices to show that $V_{info}^D - V_{contract}^D$ is increasing in κ .

Given $\lambda = 1/2$, we have

$$\frac{\partial(V_{info}^D - V_{contract}^D)}{\partial\kappa} = -\frac{N_1(\phi)N_2(\phi)}{2D_1(\phi)D_2(\phi)D_3(\phi)}$$

where

$$N_1(\phi) = 3 - b_h - 2b_m + (-5 + 2(b_h + b_m))\phi,$$

$$N_2(\phi) = -1 + b_h - 2b_h b_m + (6 - 3b_m + b_h(-4 + b_h + b_m))\phi + (-7 + 5b_h + 5b_m)\phi^2,$$

and

$$D_1(\phi) = -1 + b_h + \phi, \quad D_2(\phi) = 1 - b_m - (2 - b_h - b_m)\phi, \quad D_3(\phi) = (-1 + 2b_m)(-1 + \phi) + b_h(-1 + 2\phi).$$

With $\lambda = 1/2$, over the parameter regime D, we necessarily have $\phi_L < \phi < \phi_U$, where

$$\phi_L = \frac{b_h + 2b_m - 3}{2(b_h + b_m) - 5}, \quad \phi_U = 1 - \beta = \frac{1 - b_m}{2 - b_h - b_m} < \frac{1}{2}.$$

where ϕ_L comes from the lower bound $\underline{\phi}$ per Proposition D.1 at $\lambda = 1/2$, and ϕ_U equals $1 - \beta$ which is necessary for (D) to be optimal. Additionally, it is useful to recall that under Assumption 3 and

$$\lambda = \frac{1}{2}, 1 - b_h < \phi < b_h.$$

Under the above conditions, we prove $\frac{\partial(V_{info}^D - V_{contract}^D)}{\partial \kappa} > 0$.

- For denominators, clearly, $D_1 > 0$ and $D_2 > 0$ since $1 - b_h < \phi$ and $\phi < 1 - \beta$. Moreover, $D_3(\phi)$ is linearly increasing in ϕ . Thus, $D_3(\phi) \leq D_3(\phi_U) = 2 - b_h - \frac{3(1-b_h)}{2-b_h-b_m} < 0$ (since $\frac{1-b_h}{2-b_h-b_m} > \frac{1}{2}$ and $b_h > \frac{1}{2}$). The denominator is negative.
- For the numerators, note that $N_1(\phi)$ is decreasing in ϕ since $2(b_m + b_h) - 5 < 0$. It follows that $N_1(\phi) < N_1(\phi_L) = 0$. To evaluate the sign of N_2 , note that N_2 is decreasing in b_m because $\frac{\partial N_2}{\partial b_m} = -b_h(2 - \phi) - \phi(3 - 5\phi) < 0$ as $\phi < \frac{1}{2}$. Thus, $N_2 \leq \bar{N}_2 := N_2|_{b_m=b_h} = 2b_h^2(\phi - 1) + b_h(\phi(10\phi - 7) + 1) + \phi(6 - 7\phi) - 1$. Further differentiating $\bar{N}_2$ yields

$$\frac{\partial \bar{N}_2}{\partial b_h} = 4b_h(\phi - 1) + \phi(10\phi - 7) + 1 < -3 + \phi + 6\phi^2 < 0$$

where the first inequality follows as $b_h > 1 - \phi$ (by Assumption 3) and the second follows since $\phi \leq \frac{1}{2}$. Consequently, $\bar{N}_2$ is decreasing in b_h and therefore $\bar{N}_2 \leq \bar{N}_2|_{b_h=1-\phi} = -2(1 - 2\phi)(1 - 2\phi^2) < 0$. Hence, $N_2 < 0$.

Taken together, $D_1, D_2 > 0$, $D_3 < 0$, and $N_1, N_2 < 0$, and so $V_{info}^D - V_{contract}^D$ is linearly increasing in κ and $V_{info}^D > V_{contract}^D$ is κ is sufficiently large.

Regime E. Now we focus on the case where policy (E) is optimal for the information design approach. The existing (no-machine) contract parameter w_P^N is given by case (ii) of Proposition D.1. From the concavification approach of our analysis, the expected profit under policy (E) can be written as

$$\pi_{info}^E = \frac{\psi_2 - \lambda}{\psi_2} \pi_L + \frac{\lambda}{\psi_2} (A_I \psi_2 + B_I).$$

Substituting w_P^N and π_{info}^E yields the expression of V_{info}^E .

Given full machine information and the optimal contract to induce effort $e^*(1) = 1$, the resulting expected monetary profit is given by

$$\pi_{contract}^E = p_{11} \pi_H + (p_{10} + p_{01} + p_{00}) \pi_L.$$

Substituting $\pi_{contract}^E$ and w_P^E derived in Proposition D.2 yields the expression of $V_{contract}^E$.

Similar to the proof for regime D, since $V_{info}^E - V_{contract}^E$ is linear in κ , it suffices to show that

$V_{info}^E - V_{contract}^E$ is increasing in κ . Following an analogous approach, we first obtain

$$\frac{\partial(V_{info}^E - V_{contract}^E)}{\partial\kappa} = \frac{N_1 N_2}{2D_1 D_2 D_3}$$

where $D_1 = b_h - \phi$, $D_2 = 1 - b_h - (2 - b_h - b_m)\phi$, $D_3 = b_h + \phi - 2(b_h + b_m)\phi$, and

$$N_1 = 2 - b_h + (-5 + 2(b_h + b_m))\phi$$

and

$$N_2 = 2 - 2b_m + b_h(-1 - b_h + b_m) + (-7 + 5b_m + b_h(2 + b_h + b_m))\phi + (5 - b_h - b_m)\phi^2.$$

With $\lambda = 1/2$, over the parameter regime E, we have $\phi_L < \phi < \phi_U$ where

$$\phi_L = \beta = \frac{1 - b_h}{2 - b_h - b_m} \geq \frac{1}{2}, \quad \phi_U = \frac{2 - b_h}{5 - 2b_h - 2b_m}.$$

where ϕ_L is necessary for (E) to be optimal and ϕ_U is derived from the upper bound $\bar{\phi}$ per Proposition D.1 at $\lambda = 1/2$. Additionally, recall that under Assumption 3 and $\lambda = \frac{1}{2}$, $1 - b_h < \phi < b_h$. Provided these conditions, we now evaluate the sign of $\frac{\partial(V_{info}^E - V_{contract}^E)}{\partial\kappa}$.

- For denominators, clearly, $D_1 > 0$, $D_2 < 0$. Because D_3 is decreasing in ϕ , $D_3 \leq D_3|_{\phi=\frac{1}{2}} = \frac{1}{2} - b_m < 0$.
- For numerators, because N_1 is decreasing in ϕ , $N_1 > N_1|_{\phi=\phi_U} = 0$. To determine the sign of N_2 , observe that $\frac{\partial N_2}{\partial b_m} = -2 + 5\phi - \phi^2 + b_h(1 + \phi) > 6\phi - 2 > 0$ where the inequalities follow since $b_h > \phi$ by Assumption 3 and $\phi \geq \frac{1}{2}$. Thus, $N_2 \geq \underline{N}_2 = N_2|_{b_m=b_h} = 2b_h^2\phi + b_h(-2\phi^2 + 7\phi - 3) + 5\phi^2 - 7\phi + 2$. Further differentiating $\underline{N}_2$ leads to

$$\frac{\partial \underline{N}_2}{\partial b_h} = (4b_h + 7)\phi - 2\phi^2 - 3 > 2\phi^2 + 7\phi - 3 > 0$$

where the first inequality holds as $b_h > \phi$ and the second holds since $\phi \geq \frac{1}{2}$. It follows that $\underline{N}_2$ is increasing in b_h and $\underline{N}_2 > \underline{N}_2|_{b_h=\phi} = 2(2\phi - 1)(3\phi - 1) \geq 0$. Consequently, $N_2 > 0$.

Altogether, $D_1 > 0$, $D_2, D_3 < 0$ and $N_1, N_2 > 0$ imply that $V_{info}^E - V_{contract}^E$ is linearly increasing in κ , and thus $V_{info}^E > V_{contract}^E$ for a sufficiently large κ . $\square$

E The Case of $b_m \leq b_h$: Machine Inferior to Human

The following lemma characterizes the optimal effort as a function of the posterior v , which generalizes Lemma 2 in the main paper.

Lemma E.1. *For any given posterior belief $v = \Pr(X_m = 1|s_m)$, the manager's optimal prediction effort $e^*(v)$ can be characterized as follows.*

If $\hat{k} > \bar{K}$, then $e^(v) = 0$ for all v . “Case (N)”*

If $\hat{k} \leq \underline{K}$, then $e^(v) = 1$ for all v . “Case (A)”*

Otherwise, the optimal effort is posterior-dependent. There exists two thresholds $\psi_1 \leq \psi_2$ such that $e^(v)$ is given by one of the following three cases:*

(i) for $\phi \leq \min(u_0, l_1)$, $e^(v) = 1$ iff $v \leq \psi_2$;*

(ii) for $\phi \geq \max(u_0, l_1)$, $e^(v) = 1$ iff $v \geq \psi_1$;*

(iii) for $u_0 < \phi < l_1$, $e^(v) = 1$ iff $\psi_1 \leq v \leq \psi_2$.*

Moreover, for each $i = 1, 2$, provided $0 < \psi_i < 1$, ψ_i is increasing in ϕ .

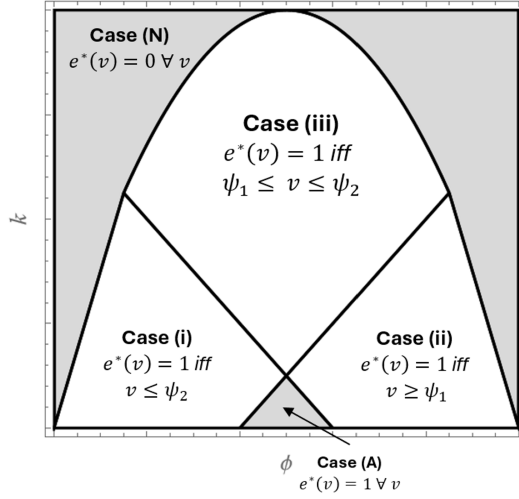


Figure E.1: Structure of the optimal effort $e^*(v)$ for different combinations of k and ϕ (Parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$)

As illustrated by Figure E.1, the structure of $e^*(v)$ is more involved than the case of $b_m > b_h$. Specifically, the two triangular regions where the manager exerts effort for $v \leq \psi_2$ and $v \geq \psi_1$ can overlap when $b_m \leq b_h$. This gives rise to a new region, Case (A), within which the manager always exert effort for any v in $[0, 1]$. In other words, in Case (A), the manager exert effort, regardless of

the messaging policy or the realized message. This happens when the effort cost $k \leq \underline{K}$. Together with the fact that the manager never exerts effort if $\hat{k} > \bar{K}$, full disclosure is trivially optimal if $\hat{k} < \underline{K}$ or $\hat{k} > \bar{K}$, since in these cases the machine prediction does influence the manager's effort.

Focusing on the parameter range of $\underline{K} \leq \hat{k} \leq \bar{K}$, we fully characterize the optimal messaging policy in Proposition E.1 below, where the closed-form expressions of η^* for different policies are given in the main paper.

Proposition E.1. *Suppose that $b_m \leq b_h$ and $\underline{K} \leq \hat{k} \leq \bar{K}$. The optimal information disclosure policies are characterized as follows.*

- (i) *Case of $\phi \leq \min(u_0, l_1)$: If $\phi \leq \beta$, full information disclosure (F) is optimal; if $\phi > \beta$, partial information disclosure is optimal with η^* given by (D') if $\lambda > \psi_2(\iff \phi < l_n)$, and by (E) if $\lambda \leq \psi_2(\iff \phi \geq l_n)$.*
- (ii) *Case of $\phi \geq \max(u_0, l_1)$: If $\phi \geq 1 - \beta$, full information disclosure (F) is optimal; if $\phi < 1 - \beta$, partial information disclosure is optimal with η^* given by (E') if $\lambda \leq \psi_1(\iff \phi \geq u_n)$, and by (D) if $\lambda > \psi_1(\iff \phi < u_n)$.*
- (iii) *Case of $u_0 < \phi < l_1$: Partial information disclosure is optimal with the cases and optimal policies defined in Table 2*

Conditions		Optimal η^*
$\phi \leq \beta$	$\lambda \geq \psi_1(\iff \phi \leq u_n)$	(D)
	$\lambda < \psi_1(\iff \phi > u_n)$	(E')
$\beta < \phi < 1 - \beta$	$\lambda > \psi_2(\iff \phi < l_n)$	(D')
	$\psi_1 < \lambda \leq \psi_2(\iff l_n \leq \phi < u_n)$	(D&E)
	$\lambda \leq \psi_1(\iff \phi \geq u_n)$	(E')
$1 - \beta \leq \phi$	$\lambda \geq \psi_2(\iff \phi \leq l_n)$	(D')
	$\lambda < \psi_2(\iff \phi > l_n)$	(E)

Table 2: Optimal η^* under case (iii) of Proposition E.1

Proposition E.1 is illustrated and discussed by Figure 8 and the accompanying discussions in the main paper. We note that in the numerical example of Figure 8, policies (E) and (D) can emerge as optimal within cases (i) and (ii) of Proposition E.1, respectively. It is possible that (D') and (E') are optimal instead, especially when the prior λ is not equal to 0.5.

F Analysis of General State-Dependent Demand Distributions

More generally, we can assume that, under each market state $\theta \in \{H, L\}$, demand d_θ follows a distribution function F_θ , where F_H is first-order stochastically larger than F_L . For any given posterior belief μ about $\theta = H$, let $\pi^*(\mu)$ denote the expected profit under the optimal solution. That is,

$$\pi^*(\mu) = \max_{Q \geq 0} r\mathbb{E}[\mu \min(d_H, Q) + (1 - \mu) \min(d_L, Q)] - cQ. \quad (\text{F.1})$$

The following lemma will be useful later in characterizing the solution structure.

Lemma F.1. $\pi^*(\mu)$ is increasing and convex in μ .

Proof of Lemma F.1. Let $\pi(\mu, Q)$ represent the expected profit given $\mu = \Pr(\theta = H)$ and any order quantity Q . Define the expected sales under state θ as $S_\theta(Q) = \mathbb{E} \min(d_\theta, Q)$. Clearly, $\frac{\partial^2 \pi(\mu, Q)}{\partial \mu \partial Q} = r(S'_H(Q) - S'_L(Q)) = r(\bar{F}_H(Q) - \bar{F}_L(Q)) \geq 0$, where $\bar{F}_\theta(x) = 1 - F_\theta(x)$. Thus, $\pi(\mu, Q)$ is supermodular in (μ, Q) , and so the optimal order quantity $Q^*(\mu)$ is increasing in μ . By the Envelope Theorem, we have

$$\frac{\partial \Pi(\mu)}{\partial \mu} = r(S_H(Q^*(\mu)) - S_L(Q^*(\mu))).$$

Clearly, $\frac{\partial \Pi(\mu)}{\partial \mu} \geq 0$ and so $\Pi(\mu)$ is increasing. Moreover, note that $S_H(Q) - S_L(Q)$ is increasing in Q , because $S'_H(Q) - S'_L(Q) = \bar{F}_H(Q) - \bar{F}_L(Q) \geq 0$ (as F_H is first-order stochastically larger). Together with $Q^*(\mu)$ increasing μ , it follows that $\frac{\partial \Pi(\mu)}{\partial \mu}$ is increasing in μ . $\square$

F.1 Uniform Demand

Our main results can be generalized analytically under the following two assumptions:

- (i) d_θ is either (a) uniformly distributed over $[a_\theta - \sigma, a_\theta + \sigma]$ with $a_H - a_L \geq 2\sigma$, or (b) uniformly distributed over $[0, a_\theta]$ with $a_H > a_L$.
- (ii) The unconditional predictions X_m and X_h are uncorrelated, i.e., $\rho(X_m, X_h) = 0$, or equivalently, $b_m + b_h = 1 + \lambda$.

Assumption (i) allows for two specifications of state-dependent uniform demand: one with disjoint supports as in (a) and the other with nested supports as in (b). Next, we show that for each of these two cases, the value of acquiring human information is unimodal in the posterior v regarding $X_m = 1$, a key property that preserves the solution structure of the base model.

F.1.1 Case (a): Uniform Demand with Disjoint Supports

Lemma F.2. *Under Assumption (i)(a) and Assumption (ii), for any posterior μ , the optimal order quantity $Q^*(\mu)$ and expected profit $\pi^*(\mu)$ are given below.*

$$Q^*(\mu) = \begin{cases} a_L + \frac{1+\mu-2\phi}{1-\mu}\sigma & \text{if } \mu \leq \phi, \\ a_H + \frac{\mu-2\phi}{\mu}\sigma & \text{if } \mu > \phi. \end{cases} \quad (\text{F.2})$$

Note that when $\mu = \frac{c}{r}$, any quantity in $[a_L + \sigma, a_H - \sigma]$ is optimal. We break the tie by choosing the smallest $Q^*(\mu) = a_L + \sigma$. Substituting $Q^*(\mu)$ into the profit function, we can derive the manager's optimal payoff as

$$\pi^*(\mu) = \begin{cases} (r - c) \left(a_L + \frac{(r\mu - c)\sigma}{r(1-\mu)} \right) & \text{if } \mu \leq \frac{c}{r}, \\ r(\mu a_H + (1 - \mu)a_L) - (a_H + \sigma)c + \frac{c^2\sigma}{r\mu} & \text{if } \mu > \frac{c}{r}. \end{cases} \quad (\text{F.3})$$

Proof of Lemma F.2. The manager's expected profit can be expressed as $r(\mu a_H + (1 - \mu)a_L) - rL(Q) - cQ$ where

$$L(Q) = \begin{cases} \mu(a_H - Q) + \frac{1-\mu}{2\sigma} \int_{Q-a_L}^{\sigma} (a_L + \epsilon - Q)d\epsilon & \text{if } Q \leq a_L + \sigma, \\ \mu(a_H - Q) & \text{if } a_L + \sigma < Q \leq a_H - \sigma, \\ \frac{\mu}{2\sigma} \int_{Q-a_H}^{\sigma} (a_H + \epsilon - Q)d\epsilon & \text{if } Q > a_H - \sigma. \end{cases} \quad (\text{F.4})$$

Since the objective is concave in Q , the optimal order quantity in (F.2) follows by the first-order optimality condition. We note that the first and second cases in (F.2) correspond to the optimal solutions attained in $[a_L - \sigma, a_L + \sigma]$ and $[a_H - \sigma, a_H + \sigma]$, respectively. (When $\mu = \frac{c}{r}$, any quantity in $[a_L + \sigma, a_H - \sigma]$ is optimal, and we break the tie by choosing the smallest $Q^*(\mu) = a_L + \sigma$.) Substituting $Q^*(\mu)$ into the profit function, we obtain the expression of the optimal expected profit as in (F.3). $\square$

With the assumption $b_m + b_h = 1 + \lambda$, that is, unconditional X_m and X_h are independent, we can simplify the belief mapping in Lemma A.1 as follows. For any posterior belief v in $X_m = 1$, we have the following posteriors regarding θ .

- If not acquiring s_h , the manager holds a posterior $\mu_N(v) = \frac{v(1-b_h)+\lambda(b_h-\lambda)}{1-\lambda}$.
- If acquiring s_h , with probability λ , the newsvender observes $s_h = 1$ which leads to $\mu(v|1) = \frac{v(1-b_h)+b_h-\lambda}{1-\lambda}$, whereas with probability $1-\lambda$, it observes $s_h = 0$ which leads to $\mu(v|0) = \frac{v(1-b_h)}{1-\lambda}$.

The value of eliciting s_h is then given by $I(v) = \lambda\pi^*(\mu(v|1)) + (1 - \lambda)\pi^*(\mu(v|0)) - \pi^*(\mu_N(v))$.

Lemma F.3. *Provided Assumption (i)(a) and Assumption (ii), the value of eliciting s_h , $I(v)$, is unimodal in $v \in [0, 1]$.*

Proof of Lemma F.3. Under the assumption of $b_m + b_h = 1 + \lambda$, $\mu_N(v)$, $\mu(v|1)$ and $\mu(v|0)$ are all linearly increasing in v . Moreover, they satisfy the following relationships.

$$\mu(v|1) = \mu_N(v) + b_h - \lambda \text{ and } \mu(v|0) = \mu_N(v) - \frac{\lambda(b_h - \lambda)}{1 - \lambda}.$$

Define $d = b_h - \lambda$ and write $\mu_N(v)$ simply as μ . Thus, $\mu(v|1) = \mu + d$ and $\mu(v|0) = \mu - \frac{\lambda d}{1 - \lambda}$. To facilitate the analysis, we change the focal variable from $v \in [0, 1]$ to $\mu \in [\frac{\lambda d}{1 - \lambda}, d]$, where $0 < d < 1 - \lambda$. Then, with abuse of notation, $I(v)$ can be expressed as a function of μ : $I(\mu) = \lambda\pi^*(\mu + d) + (1 - \lambda)\pi^*(\mu - \frac{\lambda d}{1 - \lambda}) - \pi^*(\mu)$. Because $\mu_N(v)$ is (strictly) increasing in v , to prove that $I(v)$ is unimodal, it suffices to show that $I(\mu)$ is unimodal in μ .

The expression of $I(\mu)$ is comprised of four cases, depending on the value of ϕ relative to the three possible posteriors $\mu(v|0) \leq \mu_N(v) \leq \mu(v|1)$. We will use $I_i(\mu)$ where $i = 1, 2, 3$ and 4 to denote the expression of $I(\mu)$ under the i -th case in the following analysis.

Case 1 If $\mu \leq \phi - d$, then $\mu(v|0) \leq \mu_N(v) \leq \mu(v|1) \leq \phi$. The optimal order quantity is within the low-demand support for all posteriors $\mu - \frac{\lambda d}{1 - \lambda}$, μ and $\mu + d$. As a result, $\pi^*(\cdot)$ always takes the first case in (F.3), which leads to

$$I_1(\mu) = \frac{\lambda d^2(r - c)^2\sigma/(1 - \lambda)}{r(1 - \mu)(1 - \mu - d)\left(1 - \mu + \frac{d\lambda}{1 - \lambda}\right)}.$$

Because all the factors in the denominator are positive and decreasing in μ , $I_1(\mu)$ is thus increasing.

Case 2 If $\phi - d < \mu \leq \phi$, then the order quantity is within the low-demand support given posteriors $\mu(v|0) = \mu - \frac{\lambda d}{1 - \lambda}$ and $\mu_N(v) = \mu$ but within the high-demand support given $\mu(v|1) = \mu + d$. The expression of $I_2(\mu)$ can be obtained by substituting the corresponding cases according to (F.3). We need to examine the monotonicity of $I_2(\mu)$ for $\mu \in (\phi - d, \phi]$.

First, notice that $\lim_{\mu \downarrow \phi - d} I_2'(\mu) - I_1'(\phi - d) = r\lambda(a_H - a_L - 2\sigma) \geq 0$. Because we have proved that I_1 is increasing, it follows that $\lim_{\mu \downarrow \phi - d} I_2'(\mu) > 0$.

Second, the second-order derivative of I_2 can be written as

$$I_2''(\mu) = \frac{2\sigma(r-c)^2}{r(1-\mu)^3} \left[(1-\lambda) \left(\frac{1-\mu}{1-\mu+\frac{\lambda d}{1-\lambda}} \right)^3 + \frac{\lambda c^2}{(r-c)^2} \left(\frac{1-\mu}{\mu+d} \right)^3 - 1 \right].$$

Note that the expression inside $[\cdot]$ is decreasing in μ . Therefore, $I_2''(\mu)$ must be positive, negative, or change its sign from positive to negative only once, implying that I_2 is convex, concave, or first convex and then concave in μ . Together with the fact that $I_2'(\mu) > 0$ as $\mu \rightarrow \mu - d$, over the range $(\phi - d, \phi]$, $I_2(\mu)$ must be increasing in μ if $I_2'(\phi) \geq 0$, or have a unique peak in μ if $I_2'(\phi) < 0$.

Case 3 If $\phi < \mu \leq \phi + \frac{\lambda d}{1-\lambda}$, then the order quantity is within the low-demand support given posterior $\mu(v|0) = \mu - \frac{\lambda d}{(1-\lambda)}$ but within the high-demand support for the other posteriors. According to the corresponding cases in (F.3), we can obtain the expression of $I_3(\mu)$. Then, we examine the monotonicity of I_3 for $\mu \in (\phi, \phi + \frac{\lambda d}{1-\lambda}]$.

The second-order derivative of I_3 can be written as

$$I_3''(\mu) = \frac{2\sigma c^2}{r\mu^3} \left[\lambda \left(\frac{\mu}{\mu+d} \right)^3 + \frac{(1-\lambda)(r-c)^2}{c^2} \left(\frac{\mu}{1-\mu+\frac{\lambda d}{1-\lambda}} \right)^3 - 1 \right].$$

Note that the expression inside $[\cdot]$ is increasing in μ , and so I_3'' can be negative, positive, or change its sign from negative to positive only once, implying that I_3 can be concave, convex, or first concave then convex for $\mu \in (\phi, \phi + \frac{\lambda d}{1-\lambda}]$.

As μ goes to $\phi + \frac{\lambda d}{1-\lambda}$, we have $I_3'(\phi + \frac{\lambda d}{1-\lambda}) - \lim_{\mu \downarrow \phi + \frac{\lambda d}{1-\lambda}} I_4'(\mu) = -r(1-\lambda)(a_H - a_L - 2\sigma) \leq 0$, where the expression of $I_4(\mu)$ is derived for the case of $\phi < \mu - \frac{\lambda d}{(1-\lambda)}$ and $I_4(\mu)$ is decreasing as shown below. Therefore, $I_3'(\mu) \leq 0$ at $\mu = \phi + \frac{\lambda d}{1-\lambda}$.

Then, by the concave-convex property, it follows that if $\lim_{\mu \downarrow \phi} I_3'(\mu) \leq 0$, then I_3 is decreasing throughout the range $(\phi, \phi + \frac{\lambda d}{1-\lambda}]$; if $\lim_{\mu \downarrow \phi} I_3'(\mu) > 0$, then I_3 has a unique peak within $(\phi, \phi + \frac{\lambda d}{1-\lambda}]$.

Case 4 If $\mu > \phi + \frac{\lambda d}{(1-\lambda)}$, then the order quantity is within the high-demand support for all possible posteriors, and so π^* always takes the second case in (F.3) for all posteriors. Thus, we have

$$I_4(\mu) = \frac{\lambda c^2 d^2 \sigma / (1-\lambda)}{r\mu(\mu+d) \left(\mu - \frac{\lambda d}{1-\lambda} \right)}.$$

All the factors in the denominator are positive and increasing in μ , and so $I_4(\mu)$ is decreasing.

At the boundary point $\mu = \phi$ between Cases 2 and 3, we have $\lim_{\mu \downarrow \phi} I'_3(\mu) - I'_2(\phi) = -r(a_H - a_L - 2\sigma) \leq 0$. That is, the derivative of $I(\mu)$ drops as μ increases from Case 2 to Case 3. Since $I(\mu)$ is continuous (but not differentiable at boundary points), combining the curvature of I_2 and I_3 , we can draw the following conclusion.

- If $I'_2(\phi) < 0$, then $\lim_{\mu \downarrow \phi} I'_3(\mu) < 0$. Consequently, the above case-by-case analysis implies that $I(\mu)$ is increasing in Case 1, has its unique peak within Case 2, and then keeps decreasing in Cases 3 and 4.
- If $I'_2(\phi) \geq 0$, $\lim_{\mu \downarrow \phi} I'_3(\mu)$ can be either positive or negative.
 - for $\lim_{\mu \downarrow \phi} I'_3(\mu) \geq 0$, then $I(\mu)$ is increasing in Cases 1 and 2, attains its unique peak within Case 3, and then keeps decreasing in Case 4.
 - for $\lim_{\mu \downarrow \phi} I'_3(\mu) < 0$, then $I(\mu)$ is increasing in Cases 1 and 2, attains the unique peak at $\mu = \phi$, and then becomes decreasing in Cases 3 and 4.

In all possible cases, $I(\mu)$ is unimodal in μ , and so is in v .

□

F.1.2 Case (b): Uniform Demand with Nested Supports

Lemma F.4. *Under Assumption (i)(b) and Assumption (b), for any posterior belief μ , the optimal order quantity $Q^*(\mu)$ and expected profit $\pi^*(\mu)$ are given below.*

$$Q^* = \begin{cases} \frac{a_H a_L (1-\phi)}{a_H - \mu(a_H - a_L)} & \text{if } \mu \leq \frac{\phi}{1-a_L/a_H}, \\ a_H(1 - \frac{\phi}{\mu}) & \text{if } \mu > \frac{\phi}{1-a_L/a_H}. \end{cases} \quad (\text{F.5})$$

Correspondingly, the expected monetary profit is given by

$$\pi^*(\mu) = \begin{cases} \frac{a_H a_L (r-c)^2}{2r(a_H - \mu(a_H - a_L))} & \text{if } \mu \leq \frac{\phi}{1-a_L/a_H}, \\ \frac{a_H(r\mu-c)^2 + a_L r^2 \mu(1-\mu)}{2r\mu} & \text{if } \mu > \frac{\phi}{1-a_L/a_H}. \end{cases} \quad (\text{F.6})$$

Proof of Lemma F.4. Note that the newsvendor profit as a function of the order quantity q falls into one of the two cases: If $q \in [0, a_L]$, then $\pi = r \left[\mu \int_0^q \left(1 - \frac{x}{a_H}\right) dx + (1 - \mu) \int_0^q \left(1 - \frac{x}{a_L}\right) dx \right] - cq$. If $q \in (a_L, a_H]$, then $\pi = r \left[\mu \int_0^q \left(1 - \frac{x}{a_H}\right) dx + (1 - \mu) \frac{a_L}{2} \right] - cq$. The proof is straightforward using

the first-order condition, similar to that of Lemma F.2. If $\mu \leq \frac{\phi}{1-a_L/a_H}$, the optimal order quantity falls into $[0, a_L]$; otherwise, it falls into $(a_L, a_H]$. Thus, two different sets of expressions emerge. $\square$

As before, define $I(v) = z(v)\pi^*(\mu(v|1)) + (1 - z(v))\pi^*(\mu(v|0)) - \pi^*(\mu_N(v))$ as the value of acquiring signal s_h , where $z(v)$, $\mu_N(v)$, $\mu(v|1)$, $\mu(v|0)$ are given in Lemma A.1.

Lemma F.5. *Under Assumption (i)(b) and Assumption (ii), the value of eliciting s_h , $I(v)$, is differentiable and is unimodal in $v \in [0, 1]$.*

Proof of Lemma F.5. Under the assumption of $b_m + b_h = 1 + \lambda$, $z(v) = \lambda$. Moreover, $\mu_N(v)$, $\mu(v|1)$ and $\mu(v|0)$ satisfy the following relationships.

$$\mu(v|1) = \mu_N(v) + b_h - \lambda \text{ and } \mu(v|0) = \mu_N(v) - \frac{\lambda(b_h - \lambda)}{1 - \lambda}.$$

To facilitate the analysis, define $d = b_h - \lambda$ and write $\mu_N(v)$ simply as μ . Thus, $\mu(v|1) = \mu + d$ and $\mu(v|0) = \mu - \frac{\lambda d}{1 - \lambda}$. We change the variable from $v \in [0, 1]$ to $\mu \in [\frac{\lambda d}{1 - \lambda}, d]$, where $0 < d < 1 - \lambda$. Then, with abuse of notation, $I(v)$ can be expressed as a function of μ : $I(\mu) = \lambda\pi^*(\mu + d) + (1 - \lambda)\pi^*(\mu - \frac{\lambda d}{1 - \lambda}) - \pi^*(\mu)$. Because $\mu_N(v)$ is (strictly) increasing in v , to prove that $I(v)$ is unimodal, it suffices to show that $I(\mu)$ is unimodal in μ .

Now consider the value of $I(\mu)$ for the following four cases, indexed by $i = 1, \dots, 4$.

Case 1: If $\mu(v|0) \leq \mu_N(v) \leq \mu(v|1) \leq \frac{\phi}{1-a_L/a_H}$ ($\iff \mu \leq \frac{\phi}{1-a_L/a_H} - d$), then the optimal monetary profit $\pi^*(\mu)$ always equals the first case of (F.6). Using variable μ and relationships $\mu(v|1) = \mu + d$ and $\mu(v|0) = \mu - \frac{\lambda d}{1 - \lambda}$, we obtain

$$I_1(\mu) = \frac{a_H a_L (a_H - a_L)^2 (r - c)^2 \lambda d^2}{2r(a_H(1 - \mu) + a_L \mu)(a_H(1 - \mu - d) + a_L(\mu + d))(a_H(1 - \lambda) - (a_H - a_L)(\mu(1 - \lambda) - d\lambda))}$$

for which the denominator is positive and decreasing in μ , and thus I_1 is increasing μ .

Case 2: If $\mu(v|0) \leq \mu_N(v) \leq \frac{\phi}{1-a_L/a_H} < \mu(v|1)$ ($\iff \frac{\phi}{1-a_L/a_H} - d < \mu \leq \frac{\phi}{1-a_L/a_H}$), then the optimal order quantity takes the second case of (F.5) iff the manager observes $s_h = 1$, and invoking the corresponding cases of (F.6) yields the expression of $I_2(\mu)$. One can verify that $I_2'(\mu) = I_1'(\mu) \geq 0$ at $\mu = \frac{\phi}{1-a_L/a_H} - d$. Thus, I_2 is increasing as μ starts to increase from $\frac{\phi}{1-a_L/a_H} - d$. Furthermore, define $a_H = K a_L$ where $K > 1$, we have

$$I_2''(\mu) = \frac{a_L K}{r} \left[\frac{\lambda c^2}{(\mu + d)^3} - \frac{(r - c)^2 (K - 1)^2}{(K - \mu(K - 1))^3} + \frac{(K - 1)^2 (r - c)^2 (1 - \lambda)^4}{(K(1 - \lambda) + d\lambda(K - 1) - \mu(K - 1)(1 - \lambda))^3} \right].$$

$$I_2''(\mu) = \frac{a_L(r-c)^2(K-1)^2K}{r(K-\mu(K-1))^3} \left[\frac{\lambda c^2(K-\mu(K-1))^3}{(r-c)^2(K-1)^2(\mu+d)^3} + \frac{(1-\lambda)^4(K-\mu(K-1))^3}{(K(1-\lambda)+d\lambda(K-1)-\mu(K-1)(1-\lambda))^3} - 1 \right]$$

The expression inside $[\cdot]$ is decreasing in μ . Therefore, within the interval of interest, I_2'' can change its sign at most once from positive to negative, and so I_2 must be convex, concave, or first convex and turn concave, within the interval of interest. Together with $I_2'(\mu) \geq 0$ at $\mu = \frac{\phi}{1-a_L/a_H} - d$, we can conclude that I_2 is either increasing or unimodal in $\mu \in (\frac{\phi}{1-a_L/a_H} - d, \frac{\phi}{1-a_L/a_H}]$.

Case 3: If $\mu(v|0) \leq \frac{\phi}{1-a_L/a_H} < \mu_N(v) \leq \mu(v|1)$ ($\iff \frac{\phi}{1-a_L/a_H} < \mu \leq \frac{\phi}{1-a_L/a_H} + \frac{\lambda d}{1-\lambda}$), then the optimal order quantity takes the first case of (F.5) iff the manager observes $s_h = 0$. We can thus obtain

$$I_3''(\mu) = \frac{a_L K c^2}{r \mu^3} \left[\frac{\lambda \mu^3}{(\mu+d)^3} + \frac{(K-1)^2(r-c)^2(1-\lambda)^4 \mu^3}{c^2(K(1-\lambda)+d\lambda(K-1)-\mu(K-1)(1-\lambda))^3} - 1 \right]$$

The expression inside $[\cdot]$ is increasing in μ . Thus, $I_3''(\mu)$ can be negative, positive, or change its sign from negative to positive only once as μ increases, implying that I_3 can be concave, convex, or concave-convex within the range of interest.

Case 4: If $\frac{\phi}{1-a_L/a_H} < \mu(v|0) \leq \mu_N(v) \leq \mu(v|1)$ ($\iff \mu > \frac{\phi}{1-a_L/a_H} + \frac{\lambda d}{1-\lambda}$), then the order quantity always equals the second case of (F.5), which leads to

$$I_4(\mu) = \frac{a_H c^2 d^2 \lambda}{2r \mu (\mu+d)((1-\lambda)\mu - \lambda d)}.$$

Thus, $I_4(\mu)$ is decreasing in μ . Moreover, we can verify that $I_4'(\mu) = I_3'(\mu)$ at the boundary $\mu = \frac{\phi}{1-a_L/a_H} + \frac{\lambda d}{1-\lambda}$.

With the results of Cases 3 and 4, we can conclude the following. As μ increases $\mu = \frac{\phi}{1-a_L/a_H}$ to $\mu = \frac{\phi}{1-a_L/a_H} + \frac{\lambda d}{1-\lambda}$, I_3 is eventually decreasing before it switches to Case 4. Since I_3 can be concave, convex or concave-convex, it follows that I_3 must be unimodal (or decreasing) within the range of Case 3.

Furthermore, we can verify that $I_3'(\mu) = I_2'(\mu)$ at the boundary between Cases 2 and 3, i.e., $\mu = \frac{\phi}{1-a_L/a_H}$. Together with the results of Case 2, this implies that overall $I(\mu)$ is unimodal, where the peak occurs within either Case 2 or Case 3. $\square$

F.1.3 Optimal Messaging Policy We have established that, under both specifications of uniform demand, $I(v)$ is unimodal. Because the optimal effort $e^*(v) = 1$ iff $I(v) \geq k$, our result implies that $I(v)$ and k intersect at most twice. Let $I_M = \max_{0 \leq v \leq 1} I(v)$. As long as $k < I_M$, there exists

an interval $[\psi_1, \psi_2]$ such that $e^*(v) = 1$ iff $\psi_1 \leq v \leq \psi_2$. In case of $\psi_1 < 0$, set $\psi_1 = 0$; in case of $\psi_2 > 1$, set $\psi_2 = 1$.

The expected monetary profit is given by

$$\Pi(v) = \begin{cases} \pi^*(\mu_N(v)) & \text{for } v \in [0, \psi_1) \cup (\psi_2, 1] \\ \lambda\pi^*(\mu(v|1)) + (1 - \lambda)\pi^*(\mu(v|0)) & \text{for } v \in [\psi_1, \psi_2] \end{cases}$$

By Lemma F.1, $\Pi(v)$ consists of (at most) three *convex* pieces because convexity is preserved under linear composition with $\mu_N(v)$, $\mu(v|1)$ and $\mu(v|0)$, where the linearity of $\mu(v|1)$ and $\mu(v|0)$ is due to Assumption (ii). Therefore, the upper concave envelope of $\Pi(v)$ is constructed in the same fashion as in our base model with points $(0, \Pi(0))$, $(\psi_1, \Pi(\psi_1))$, $(\psi_2, \Pi(\psi_2))$, $(1, \Pi(1))$.

Based on the above observations, we can write out the conditions for each possible concave envelope and hence the corresponding optimal η^* , as summarized in Proposition 5 in the main paper. Note that the conditions in Proposition 5 also accommodate the cases where $\psi_1 = 0$ or $\psi_2 = 1$.

It is worth noting that the expected monetary profit $\Pi(v)$ is convex within each case, as opposed to linear in our base model. Intuitively, the convexity arises because, unlike the base model where human forecast provides value only when the realized s_h influences the order choice between a_H and a_L , now a marginal amount of information improves the expected profit because the optimal order quantity is continuously adjusted. Consequently, whenever (D&E) is optimal, it is also the unique optimal messaging policy.

F.2 Normal Demand: Numerical Study

When d_θ follows a more general distribution or the unconditional X_m and X_h has a correlation coefficient $\rho(X_m, X_h) > 0$, unfortunately, it appears intractable to generalize our analytical results. To test the robustness of our main results, we conduct a numerical study. We consider the case where d_θ follows a Normal distribution $N(a_\theta, \sigma^2)$ for each $\theta \in \{H, L\}$ and $a_H > a_L$. We discretize the support of v into 1001 points and for each $v = 0, .001, .002, \dots, 1$, we evaluate $\Pi(v)$; then we construct the upper concave envelope over the 1001 points. This exhaustive approach enables us to verify if the optimal messaging policy aligns with one of the forms characterized in the base model. Figure F.1 reports the optimal policies for various combinations of k and ϕ .

In all instances, the optimal policy always takes one of the forms introduced in the main paper. The case of $\beta = 0.55$ in Figure F.1(a) exhibits a similar structure as in Proposition 2, except for

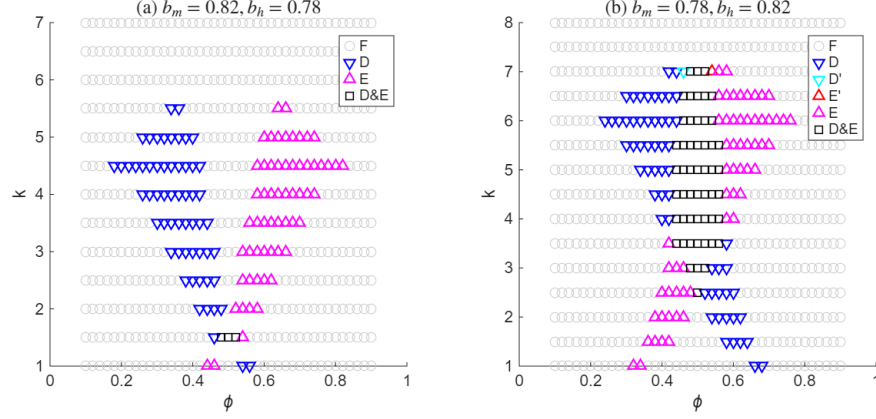


Figure F.1: Structure of the optimal messaging policies with $d_\theta \sim N(a_\theta, 3^2)$ [parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$ and c is varied according to ϕ given in the graph]

very small k ; for example, when $k = 1.5$, we observe (D&E) can be optimal. The case of $\beta = 0.45$ in Figure F.1(b) is in line with the solution structure for the case of $b_m \leq b_h$, as discussed in §5.2.

To explain the reason why (D&E) can be optimal for $\beta > 1/2$, Figure F.2 presents several representative instances when $\beta = 0.55$ and the manager would never exert effort under full disclosure, corresponding to Case (iii) of the base model. In line with our base model, when $k = 5$, (D) and (E) are optimal for small and high ϕ , respectively, whereas full disclosure is optimal for $\phi = 0.5$. When $k = 1.5$, (D&E) is optimal for $\phi = 0.5$, instead of (F). Although full machine information is better than no machine information (i.e., at $v = \lambda = 0.5$, full-information profit is higher than $\Pi(0.5)$), the upper concave envelope of $\Pi(v)$ can hinge on $v = \psi_1$ and $v = \psi_2$ such that (D&E) is optimal, because $\Pi(v)$ is convex in $v \in [\psi_1, \psi_2]$ rather than linear. Also, because of this, policy (D&E) is strictly better than providing no machine information.

G Analysis of the Continuous-Effort Model

We allow the manager to choose a continuous effort level $e \in [0, 1]$ while incurring a convex increasing cost ke^2 where $k > 0$. We follow the “truth-or-noise” setting widely adopted in the literature (Johnson and Myatt, 2006; Hagiu and Wright, 2020). With probability e , s_h reveals the true value of X_h , i.e., $s_h = X_h$. With probability $1 - e$, s_h is an independent draw from the same distribution as X_h , and thus not informative at all. Consequently, s_h and X_h are correlated according to the following conditional probabilities: $\Pr(X_h = 1|s_h = 1) = e + (1 - e)\lambda$ and $\Pr(X_h = 0|s_h = 0) = e + (1 - e)(1 - \lambda)$. This effort-dependent prediction is also in line with Taylor and Xiao (2009). By the law of total probability, it is easy to verify that $\Pr(\theta = H|s_h = H) = b_h e + (1 - e)\lambda$ and

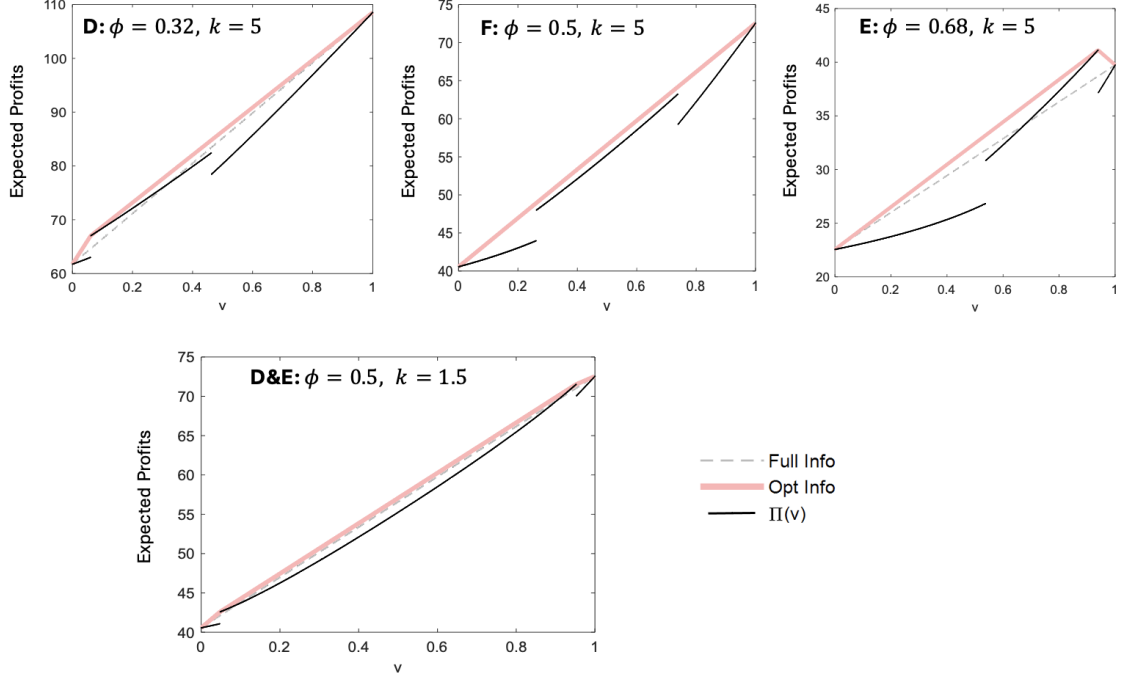


Figure F.2: Expected profit functions in representative instances for the case of $\beta = 0.55$ in Figure F.1

$\Pr(\theta = L|s_h = L) = b'_h e + (1 - e)(1 - \lambda)$. In this regard, we can interpret X_h as the best prediction that a human can generate.

In the following, we first present the analytical results in §G.1, followed by numerical illustrations in §G.2. The technical proofs are provided in §G.3

G.1 Analytical Results

The manager's posterior belief regarding $\theta = H$ is generally a function of effort level e , posterior $v = \Pr(X_m = 1|s_m)$, and the realized signal s_h , denoted by $\mu(e, v|s_h)$. In the proof of Lemma A.1, we have obtained the general expressions of $\mu(e, v|1)$ and $\mu(e, v|0)$. Recall that $\mu(e, v|s_h)$ where $s_h \in \{0, 1\}$ is the posterior probability of $\theta = H$ given any effort level $e \in [0, 1]$, posterior belief $\Pr(X_m = 1) = v$, and realized signal s_h . For ease of reference, their expressions are provided below.

$$\mu(v, e|1) = \frac{e(1 - \lambda)[v(b_h + b_m - 1 - \lambda) + \lambda(1 - b_m)] + \lambda(\lambda(1 - b_m) + v(b_m - \lambda))}{\lambda(1 - \lambda) + e(b_m + b_h - 1 - \lambda)(v - \lambda)}. \quad (\text{G.1})$$

and

$$\mu(v, e|0) = \frac{(1 - \lambda)[\lambda(1 - b_m) + v(b_m - \lambda) - e(v(b_m + b_h - 1 - \lambda) + \lambda(1 - b_m))]}{(1 - \lambda)^2 - e(b_m + b_h - 1 - \lambda)(v - \lambda)}. \quad (\text{G.2})$$

To begin, we note some useful properties of $\mu(e, v|1)$ and $\mu(e, v|0)$.

Lemma G.1. *For any $e \in [0, 1]$, the posterior belief $\mu(e, v|s_h)$ given realized $s_h \in \{0, 1\}$ satisfies the following properties: (i) $\mu(e, v|0) \leq \mu(e, v|1)$ where the equality holds at $e = 0$; (ii) $\mu(e, v|1)$ is increasing in e whereas $\mu(e, v|0)$ is decreasing in e .*

As implied by the above, $\mu(e, v|0) \leq \mu(e, v|1)$ where the equality holds at $e = 0$. Intuitively, a greater effort amplifies the distinction between $\mu(e, v|1)$ and $\mu(e, v|0)$, making s_h more informative.

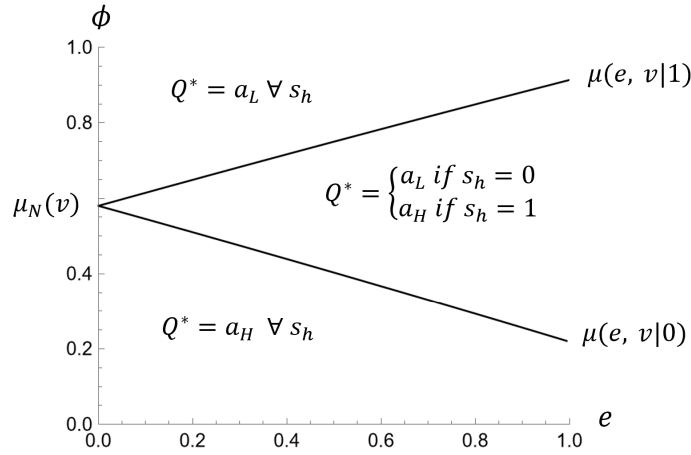


Figure G.1: Illustration of posterior beliefs $\mu(e, v|1)$ and $\mu(e, v|0)$ as functions of e with v fixed.

Figure G.1 illustrates how $\mu(e, v|1)$ and $\mu(e, v|0)$, resulting from any effort $e \in [0, 1]$, affects the optimal order quantity. By Lemma 1, the optimal order quantity will be contingent on the realization of s_h if $\mu(e, v|0) \leq \phi < \mu(e, v|1)$; see the triangular region between $\mu(e, v|0)$ and $\mu(e, v|1)$ in Figure G.1. Yet, the optimal order quantity will not be impacted by s_h if $\phi < \mu(e, v|0)$ or $\phi \geq \mu(e, v|1)$. Thus, choosing a positive effort $e > 0$ brings value only if e is sufficiently large such that s_h is informative enough, resulting in $\mu(e, v|0) \leq \phi < \mu(e, v|1)$.

Given any posterior $v = \Pr(X_m = 1)$, the human chooses an effort level $e \in [0, 1]$ to maximize the following expected payoff:

$$\begin{aligned} u(e, v) &= \mathbb{E}_{s_h \sim z(e, v)} [\pi^*(\mu(e, v|s_h))] - ke^2 \\ &= z(e, v)\pi^*(\mu(e, v|1)) + (1 - z(e, v))\pi^*(\mu(e, v|0)) - ke^2. \end{aligned} \quad (\text{G.3})$$

where function π^* is as given in Lemma 1 and $z(e, v)$ denotes the probability of $s_h = 1$ which is given by $z(e, v) = \frac{1}{1-\lambda}[\lambda(1-\lambda) + e(v-\lambda)(b_m + b_h - 1 - \lambda)]$.

As discussed earlier, the information elicitation effort improves the expected monetary profit only if $\mu(e, v|0) \leq \phi < \mu(e, v|1)$ such that the realization of s_h indeed influences the order quantity. Consequently, it is optimal to either exert zero effort or set the effort level sufficiently high such that $\mu(e, v|0) \leq \phi < \mu(e, v|1)$.

Define $u^i(e, v) = (r - c)a_L$, $u^{ii}(e, v) = (r - c)a_L + (r\mu_N(v) - c)(a_H - a_L) - ke^2$, and $u^{iii}(e, v) = (r - c)a_L + z(e, v)(r\mu(e, v|1) - c)(a_H - a_L) - ke^2$. Note that u^i is the manager's payoff when order quantity $Q = a_L$, regardless of signal s_h , u^{ii} is that when $Q = a_H$, regardless of signal s_h , and u^{iii} is that when $Q = a_L$ given $s_h = 0$ but $Q = a_H$ given $s_h = 1$.

Further, define $U_N(v) = \max(u^i(0, v), u^{ii}(0, v))$ as the payoff function given $e = 0$. That is, $U_N(v)$ is the maximum payoff with any information elicitation effort, in which case Q is independent of s_h . Define $U_I(v) = u^{iii}(e^*(v), v)$ where $e^*(v) = \arg \max_{e \in [0, 1]} u^{iii}(e, v)$. That is, $U_I(v)$ is the maximum payoff with information elicitation, assuming the order quantity to be adjusted according to the realization of s_h .

The following lemma establishes that the optimal effort level can be determined by comparing U_I and U_N .

Lemma G.2. *The optimal effort level $e^*(v)$ is determined by the following rule: If $U_N(v) \geq U_I(v)$, then $e^*(v) = 0$; otherwise, $e^*(v) = \min(e^o(v), 1)$ where $e^o(v) = \frac{r(a_H - a_L)h(v)}{k}$ and*

$$h(v) = \frac{[v(1 - \lambda - \phi) + \lambda\phi](b_m + b_h - 1 - \lambda) + \lambda(1 - \lambda)(1 - b_m)}{2(1 - \lambda)}. \quad (\text{G.4})$$

For technical convenience, hereafter we impose a sufficient condition that the relative effort cost $\hat{k} = \frac{k}{r(a_H - a_L)}$ is large enough, which ensures that the optimal effort level is either zero or an interior solution $e^o(v)$ less than one that maximizes $u^{iii}(e, v)$, and that the interval $[\psi_1, \psi_2]$ lies strictly within $[0, 1]$.

Proposition G.1. *Suppose that $\hat{k} = \frac{k}{r(a_H - a_L)} > \underline{K}$. There exists two thresholds $\psi_1, \psi_2 \in (0, 1)$ such that the optimal effort level $e^*(v) = 0$ if $v < \psi_1$ or $v > \psi_2$, and $e^*(v) = e^o(v)$ if $\psi_1 \leq v \leq \psi_2$, where $e^o(v) = h(v)/\hat{k}$, ψ_1 is the larger root solving the quadratic equation $U^I(v) = (r - c)a_L$, and ψ_2 is the smaller root solving a quadratic equation $U^I(v) = (r - c)a_L + (r\mu_N(v) - c)(a_H - a_L)$.*

Although the expressions of ψ_1 and ψ_2 become more complex, they can be determined by solving

two quadratic equations as stated in Proposition G.1. Using ψ_1 and ψ_2 , we can derive conditions under which partial disclosure is optimal, taking the same forms as defined in our base model.

Proposition G.2. *Suppose that $\hat{k} = \frac{k}{r(a_H - a_L)} > \underline{K}$. The optimal messaging policy is characterized as follows. If $\frac{h^2(\psi_1)}{\hat{k}} \leq \psi_1(b_m - \phi)$ and $\frac{h^2(\psi_2)}{\hat{k}} \leq \psi_2 b_m + (1 - \psi_2)\phi - \mu_N(\psi_2)$, full disclosure is optimal. Otherwise, partial disclosure is optimal with the following cases:*

- (1) *If $\frac{h^2(\psi_1)}{\hat{k}} > \psi_1(b_m - \phi)$ and $\frac{1}{\hat{k}} \left(\frac{1 - \psi_1}{1 - \psi_2} h^2(\psi_2) - h^2(\psi_1) \right) \leq \frac{\psi_2 - \psi_1}{1 - \psi_2} (b_m - \phi) - \frac{1 - \psi_1}{1 - \psi_2} \mu_N(\psi_2)$, then η^* is given by (E') for $\lambda \leq \psi_1$, and by (D) for $\lambda > \psi_1$;*
- (2) *if $\frac{h^2(\psi_2)}{\hat{k}} > \psi_2 b_m + (1 - \psi_2)\phi - \mu_N(\psi_2)$ and $\frac{1}{\hat{k}} \left(h^2(\psi_2) - \frac{\psi_2}{\psi_1} h^2(\psi_1) \right) \geq \phi - \mu_N(\psi_2)$, η^* is given by (E) for $\lambda \leq \psi_2$, and by (D') for $\lambda > \psi_2$;*
- (3) *otherwise, η^* is given by (E') for $\lambda \leq \psi_1$, by (D&E) for $\psi_1 < \lambda \leq \psi_2$, and (D') for $\lambda > \psi_2$.*

G.2 Numerical Illustration

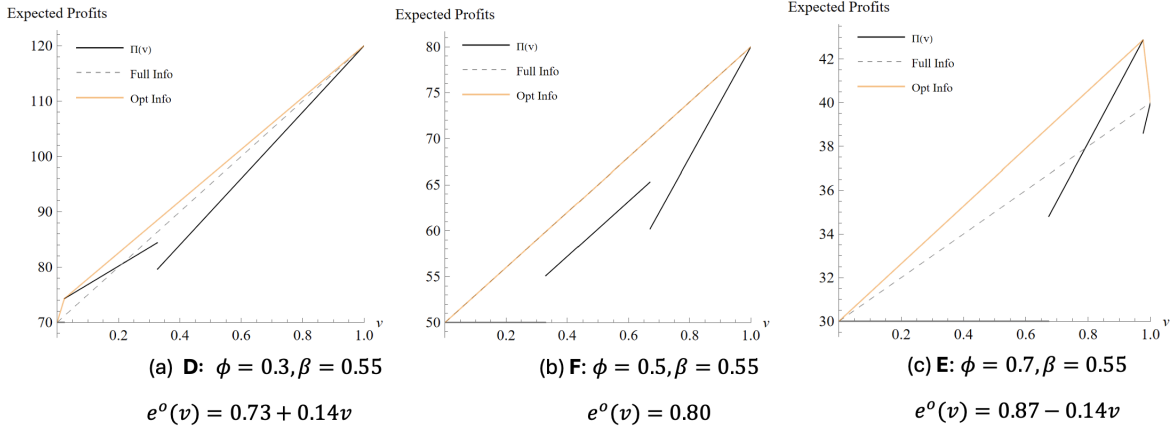


Figure G.2: Illustration of the optimal messaging policies (D), (F) and (E) with continuous effort [parameter values: $\lambda = 0.5$, $a_H = 20$, $a_L = 10$, $r = 10$, $k = 8$, $b_m = 0.8$ whereas b_h and c are set according to the values of β and ϕ given in the graph.]

Figure G.2 present numerical examples in which, given $\beta > 1/2$, partial disclosure policies are optimal for $\phi = 0.3$ and $\phi = 0.7$, respectively, whereas full information is optimal for $\phi = 0.5$, similar to our results discussed in §4.3.2. Note that the optimal effort $e^*(v)$ is less than one and can be a function of v . Moreover, Figure G.3 provides an example where (D&E) is optimal given $\beta < 1/2$. In the example, $\lambda = 0.6$ and so garbling *both* message realizations is optimal. Notably, different from our base model, we can show that $\Pi(v)$ is convex for $v \in [\psi_1, \psi_2]$, rather than linear. Consequently, (D&E) is the unique optimal policy, strictly outperforming the case of providing

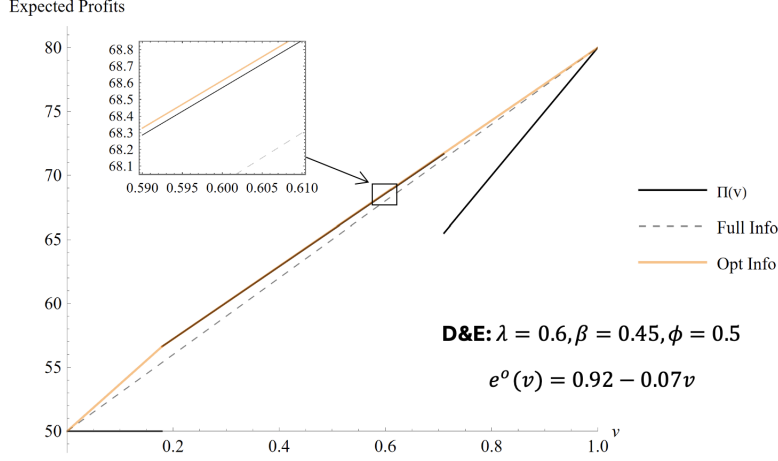


Figure G.3: Illustration of the optimal messaging policy (D&E) with continuous effort [parameter values: $\lambda = 0.6$, $a_H = 20$, $a_L = 10$, $r = 10$, $k = 8$, $b_m = 0.8$ whereas b_h and c are set according to the values of β and ϕ given in the graph.]

no machine information. As demonstrated in the magnified graph, the optimal monetary profit is strictly higher than $\Pi(0.6)$, i.e., the profit based on the prior $\lambda = 0.6$.

G.3 Proofs

Proof of Lemma G.1. First, notice that at $e = 0$, $\mu(e, v|1) = \mu(e, v|0) = \mu_N(v)$ where

$$\mu_N(v) = \frac{v(b_m - \lambda) + \lambda(1 - b_m)}{1 - \lambda}.$$

By examining the first-order derivatives, one can show that $\frac{\partial \mu(e, v|1)}{\partial e} \geq 0$ and $\frac{\partial \mu(e, v|1)}{\partial e} \leq 0$ for any $v \in [0, 1]$ within the assumed parameter range where $b_m + b_h \geq 1 + \lambda$. Thus, property (ii) follows. $\square$

Proof of Lemma G.2. To maximize $u(e, v)$ in (G.3) over e , we consider three cases.

- (i) If $\phi \geq \mu(e, v|1)$, regardless of s_h , the manager orders a_L units, leading to $\pi^*(\mu(e, v|0)) = \pi^*(\mu(e, v|1)) = (r - c)a_L$. Consequently, the manager's payoff equals $u(e, v) = u^i(e, v) = (r - c)a_L - ke^2$, decreasing in e .
- (ii) If $\phi < \mu(e, v|0)$, regardless of s_h , the manager orders a_H units, leading to $u(e, v) = (r - c)a_L + (r\mathbb{E}_{s_h \sim z(e, v)}[\mu(e, v|s_h)] - c)(a_H - a_L) - ke^2$, where the posterior μ unconditional on s_h ,

$$\mathbb{E}_{s_h \sim z(e, v)}[\mu(e, v|s_h)] = \mu_N(v) = \frac{v(b_m - \lambda) + \lambda(1 - b_m)}{1 - \lambda}$$

which is independent of e . Thus, the manager's payoff is given by $u(e, v) = u^{ii}(e, v) = (r - c)a_L + (r\mu_N(v) - c)(a_H - a_L) - ke^2$, decreasing in e as well.

(iii) If $\mu(e, v|0) \leq \phi < \mu(e, v|1)$, then the optimal order quantity is a_L given $s_h = 0$ but is a_H given $s_h = 1$. As a result, the payoff $u(e, v)$ is equal to

$$u^{iii}(e, v) = (r - c)a_L + z(e, v)(r\mu(e, v|1) - c)(a_H - a_L) - ke^2$$

where the second term is linear in e . Hence, the payoff $u^{iii}(e, v)$ is quadratic concave in e . Moreover, the unconstrained maximizer of u^{iii} is given by

$$e^o(v) = \frac{r(a_H - a_L)h(v)}{k} \quad (\text{G.5})$$

where

$$h(v) = \frac{[v(1 - \lambda - \phi) + \lambda\phi](b_m + b_h - 1 - \lambda) + \lambda(1 - \lambda)(1 - b_m)}{2(1 - \lambda)}. \quad (\text{G.6})$$

Note that $h(v)$ is linear in v , and that $h(0) \geq 0$ and $h(1) \geq 0$. Thus, $e^o(v) \geq 0$ for all $v \in [0, 1]$. In general, the optimal effort e maximizing u^{iii} is given by $e^I = \min(e^o, 1)$, where the minimum operator $\min(e^o, 1)$ incorporates the possibility of $e^o > 1$ in which case the maximizer over $[\hat{e}_1, 1]$ is $e = 1$.

By the fact that $\mu(0, v|1) = \mu(1, v|1) = \mu_N(v)$, $\mu(e, v|1)$ is increasing and $\mu(e, v|0)$ is decreasing in e , we have two different cases when optimizing over $e \in [0, 1]$. Note that

$$\mu_N(v) \leq \phi \iff v \leq v_N := \frac{(1 - \lambda)\phi - \lambda(1 - b_m)}{b_m - \lambda}.$$

Case of $v \leq v_N$, i.e., $\mu_N(v) \leq \phi$. Let $\hat{e}_1$ be the unique solution such that $\mu(\hat{e}_1, v|1) = \phi$. Then, $e \leq \hat{e}_1$ iff $\mu(e, v|1) \leq \phi$.

In this case, $u(e, v)$ consists of two pieces. For $e \leq \hat{e}_1$, $u(e, v) = u^i(e, v)$ which is decreasing in e ; for $e > \hat{e}_1$, $u(e, v) = u^{iii}(e, v)$ which is concave in e . For any $e \in [0, \hat{e}_1]$, $u(e, v)$ is maximized at $e = 0$, resulting in the maximum payoff $u^i(0, v) = (r - c)a_L$.

Recall that we use e^I to denote the maximizer of u^{iii} . In the following, we argue that the optimal effort equals $e^* = e^I$ if $u^{iii}(e^I, v) \geq u^i(0, v)$, and $e^* = 0$ otherwise. To see this, if $e^I > \hat{e}_1$, then the optimal effort is determined exactly by comparing $u^{iii}(e^I, v)$ and $u^i(0, v)$. If $e^I \leq \hat{e}_1$, the optimal effort $e^* = 0$, and it must be the case that $u^{iii}(e^I, v) \leq u^i(0, v)$. This is because

$u^{iii}(e^I, v) \leq u^i(e^I, v) \leq u^i(0, v)$, where the first inequality follows because, provided $e^I \leq \hat{e}_1$ ($\iff \mu(e^I, v|1) \leq \phi$), ordering a_L units is optimal under both signal realizations of s_h (and hence $u^i > u^{iii}$ for any given e), and the second inequality follows since u^i is decreasing in e .

Case of $v > v_N$, i.e., $\mu_N(v) > \phi$. Let $\hat{e}_0$ be the unique solution such that $\mu(\hat{e}_0, v|0) = \phi$. Then, $e \leq \hat{e}_0$ iff $\mu(e, v|0) \geq \phi$. Accordingly, $u(e, v)$ consists of two pieces. For $e \leq \hat{e}_0$, $u(e, v) = u^{ii}(e, v)$ which is decreasing in e ; for $e > \hat{e}_0$, $u(e, v) = u^{iii}(e, v)$ which is concave in e . Analogous to the previous case, the optimal effort level $e^*(v)$ is either 0 or e^I , depending on whether $u^{iii}(e^I, v)$ is greater than $u^{ii}(0, v)$. The reason is similar. If $e^I > \hat{e}_0$, $e^*(v)$ is determined exactly by comparing $u^{ii}(0, v)$ and $u^{iii}(e^I, v)$. If $e^I \leq \hat{e}_0$, we must have $u^{iii}(e^I, v) \leq u^{ii}(e^I, v) \leq u^{ii}(0, v)$ where the first inequality follows because, provided $e^I \leq \hat{e}_0$ ($\iff \mu(e^I, v|0) \geq \phi$), ordering a_H units regardless of the realization of s_h is optimal and hence u^{ii} is higher than u^{iii} for any given e , and the second inequality is due to the fact that u^{ii} is decreasing in e .

$$\text{Note that } u(0, v) = \begin{cases} u^i(0, v) & \text{if } v \leq v_N \\ u^{ii}(0, v) & \text{if } v > v_N \end{cases} \text{ can be written as } u(0, v) = \max(u^i(0, v), u^{ii}(0, v))$$

We can therefore put together the above two cases and conclude that to determine the optimal $e^*(v)$, it suffices to compare $U_N(v) = u(0, v)$ and $U_I(v) = u^{iii}(e^I, v)$. $\square$

Proof of Proposition G.1. For technical convenience, we assume that k is large enough such that $e^o(v) \leq 1$. By the expression of $e^o(v)$, a sufficient condition is

$$\frac{k}{r(a_H - a_L)} \geq K_{interior} := \frac{1}{2} \left(\lambda(1 - b_m) + (b_m + b_h - 1 - \lambda) \max\left(\frac{\lambda\phi}{1 - \lambda}, 1 - \phi\right) \right) \quad (\text{G.7})$$

under which $e^o(v) \leq 1$ for all $v \in [0, 1]$. We first establish the following lemma.

Lemma G.3. *Under condition (G.7), we have (i) $U_I(v)$ is convex and increasing in v ; (ii) $U_I'(0) > \frac{\partial u^i(0, v)}{\partial v}$, and $U_I'(1) < \frac{\partial u^{ii}(0, v)}{\partial v}$.*

Proof of Lemma G.3. For part (i), note that $U_I(v) = u^{iii}(e^o, v)$ is a quadratic function of v . It is straightforward to verify that $U_I''(v) = \frac{(a_H - a_L)^2(r(1 - \lambda) - c)^2(b_m + b_h - \lambda - 1)^2}{2k(1 - \lambda)^2} \geq 0$, and thus U_I is convex.

To show that $U_I(v)$ is increasing, we apply the Envelope Theorem, $U_I'(v) = \frac{\partial u^{iii}(e, v)}{\partial v}|_{e=e^o}$. Notice that $\frac{\partial u^{iii}(e, v)}{\partial v}$ is linear in e , $\frac{\partial u^{iii}(e, v)}{\partial v}|_{e=0} = \frac{r\lambda(a_H - a_L)(b_m - \lambda)}{1 - \lambda} \geq 0$, and $\frac{\partial u^{iii}(e, v)}{\partial v}|_{e=1} = \frac{(a_H - a_L)[r(b_m + b_h - b_h\lambda - 1) - c(b_m + b_h - \lambda)]}{1 - \lambda}$ (because $b_h \leq 1$ and $r > c$). It follows that $\frac{\partial u^{iii}(e, v)}{\partial v} \geq 0$ for any $e \in [0, 1]$, and so $U_I'(v) = \frac{\partial u^{iii}(e, v)}{\partial v}|_{e=e^o} \geq 0$.

For part (ii), clearly, $U_I'(0) - \frac{\partial u^i(0, v)}{\partial v} \geq 0$ as $u^i(0, v) = (r - c)a_L$ is constant in v . One can verify that, at $v = 1$, $\frac{\partial u^{ii}(0, v)}{\partial v} - \frac{\partial u^{iii}(e, v)}{\partial v} = \frac{a_H - a_L}{1 - \lambda} T(e)$ where $T(e)$ is linear in e with $T(0) =$

$r(b_m - \lambda)(1 - \lambda) > 0$ and $T(1) = c(b_m + b_h - 1 - \lambda) + r(1 - \lambda)(1 - b_h) > 0$. Thus, $\frac{\partial u^{ii}(0,v)}{\partial v} > \frac{\partial u^{iii}(e,v)}{\partial v}$ for any $e \in [0, 1]$. Due to the Envelope Theorem, we have $\frac{\partial u^{ii}(0,v)}{\partial v} > U'_I(1)$. $\square$

Note that $U_N(v)$ is a piece-wise linear function with the first piece having a slope $\frac{\partial u^i(0,v)}{\partial v} = 0$ for $v < v_N$ and the second piece having a slope $\frac{\partial u^{ii}(0,v)}{\partial v}$ for $v \geq v_N$. By Lemma G.3, it follows that for $v \in [0, 1]$, U_I can intersect with the first piece of U_N at most once, and can cross the second piece of U_N also at most once. Moreover, at the inflection point $v = v_N$, we have $U_I(v_N) \geq U_N(v_N) = (r - c)a_L$ because

$$U_I(v_N) - U_N(v_N) = \frac{(a_H - a_L)^2((1 - b_m)(1 - b_h)r^2\lambda + c(r - c)(b_m + b_h - 1 - \lambda))^2}{4kr^2(b_m - \lambda)^2} > 0.$$

Consequently, U_I must intersect with the first piece of U_N at some point $v = \psi_1 \in [0, v_N]$, provided $U_I(0) < u^i(0, v)$, which is equivalent to

$$K_1 := \frac{\lambda[(1 - \lambda)(1 - b_m) + \phi(b_m + b_h - 1 - \lambda)]^2}{4(1 - \lambda)[(1 - \lambda)\phi - \lambda(1 - b_m)]} < \frac{k}{r(a_H - a_L)} \quad (\text{G.8})$$

Similarly, U_I must intersect with the second piece of U_N at some point $v = \psi_2 \in [v_N, 1]$, provided $U_I(1) < u^{ii}(0, 1)$, or equivalently,

$$K_2 := \frac{[\lambda(1 - b_m) + (1 - \phi)(b_m + b_h - 1 - \lambda)]^2}{4(1 - \lambda)(b_m - \phi)} < \frac{k}{r(a_H - a_L)}. \quad (\text{G.9})$$

Moreover, by the convexity of U_I , ψ_1 is the larger root solving the quadratic equation $U_I(v) = u^i(0, v)$, and ψ_2 is the smaller root solving $U_I(v) = u^{ii}(0, v)$.

The conditions assumed in the above analysis are (G.7), (G.8) and (G.9). They can be summarized as $\frac{k}{r(a_H - a_L)} > \underline{K}$ where $\underline{K} = \max(K_{interior}, K_1, K_2)$. $\square$

Proof of Proposition G.2. Now consider the firm's problem to maximize the expected monetary profit, which can be expressed as a function consisting of three cases.

$$\Pi(v) = \begin{cases} u^i(0, v) & \text{for } v < \psi_1, \\ U_I(v) + ke^o(v)^2 & \text{for } \psi_1 \leq v \leq \psi_2, \\ u^{ii}(0, v) & \text{for } v > \psi_2. \end{cases}$$

Thus, $\Pi(v)$ consists of three pieces: (1) the first piece $u^i(0, v) = (r - c)a_L$ is constant in v ; (2) the

second piece equals $U_I(v) + ke^o(v)^2$ as the effort cost is not internalized in the monetary profit, and moreover, $U_I(v) + ke^o(v)^2$ is convex in v ; (3) the third piece $u^{iii}(0, v) = (r-c)a_L + (r\mu_N v - c)(a_H - a_L)$ is linearly increasing in v . Additionally, $\Pi(v)$ is discontinuous at the two breakpoints $v = \psi_1$ and ψ_2 .

Let $y_0 = (r-c)a_L$, $y_{\psi_1} = (r-c)a_L + ke^o(\psi_1)^2$, $y_{\psi_2} = u^{ii}(0, \psi_2) + ke^o(\psi_2)^2 = (r-c)a_L + (r\mu_N(\psi_2) - c)(a_H - a_L) + ke^o(\psi_2)^2$, and $y_1 = u^{ii}(0, 1) = (r-c)a_L + (r\mu_N(1) - c)(a_H - a_L)$. Because $\Pi(v)$ consists of three pieces that are either linear or convex in v , the upper concave envelope of $\Pi(v)$ is determined by the positions of four points: $(0, y_0)$, (ψ_1, y_{ψ_1}) , (ψ_2, y_{ψ_2}) , and $(1, y_1)$. Comparing the slopes of different line segments between these points, we can derive conditions for each of the information provision policies to be optimal.

The following four conditions are relevant to determine the upper concave envelope.

$$\frac{y_{\psi_1} - y_0}{\psi_1} > y_1 - y_0 \iff \frac{r(a_H - a_L)}{k} h^2(\psi_1) > \psi_1(b_m - \phi), \quad (\text{f1})$$

$$\begin{aligned} \frac{y_1 - y_{\psi_1}}{1 - \psi_1} &\leq \frac{y_1 - y_{\psi_2}}{1 - \psi_2} \iff \\ \frac{r(a_H - a_L)}{k} \left(\frac{1 - \psi_1}{1 - \psi_2} h^2(\psi_2) - h^2(\psi_1) \right) &\leq \frac{\psi_2 - \psi_1}{1 - \psi_2} (b_m - \phi) - \frac{1 - \psi_1}{1 - \psi_2} \mu_N(\psi_2), \end{aligned} \quad (\text{m1})$$

$$\frac{y_{\psi_2} - y_0}{\psi_2} > y_1 - y_0 \iff \frac{r(a_H - a_L)}{k} h^2(\psi_2) > \psi_2 b_m + (1 - \psi_2)\phi - \mu_N(\psi_2), \quad (\text{f2})$$

$$\frac{y_{\psi_2} - y_0}{\psi_2} \geq \frac{y_{\psi_1} - y_0}{\psi_1} \iff \frac{r(a_H - a_L)}{k} \left(h^2(\psi_2) - \frac{\psi_2}{\psi_1} h^2(\psi_1) \right) \geq \phi - \mu_N(\psi_2). \quad (\text{m2})$$

If neither (f1) nor (f2) holds, then the upper concave envelop is a linear line segment connecting two points $(0, y_0)$ and $(1, y_1)$, implying that full information provision is optimal. Otherwise, if (f1) or (f2) holds true, partial information provision is optimal. Specifically, (i) if (f1) and (m1) hold, the upper concave envelop is piece-wise linear connecting three points $(0, y_0)$, (ψ_1, y_{ψ_1}) , and $(1, y_1)$. (ii) if (f2) and (m2) hold, the upper concave envelop is piece-wise linear connecting three points $(0, y_0)$, (ψ_2, y_{ψ_2}) , and $(1, y_1)$. (iii) otherwise, the upper concave envelop is piece-wise linear consisting of all the four points. Using the same graphic approach as in our base model, we can conclude the various cases stated in Proposition G.2. $\square$