

A Reputation-Based Contract for Repeated Crowdsensing With Costly Verification

Donya Ghavidel Dobakhshari , Parinaz Naghizadeh , Mingyan Liu , and Vijay Gupta 

Abstract—A system operator asks a group of sensors to exert costly effort to collect accurate measurements of a value of interest over time. At each time, each sensor is asked to report its observation to the operator, and is suitably compensated for the costly effort it exerts. Since both the effort and the observation are private information for the sensor, a naive payment scheme which compensates the sensor based only on its self-reported values of the effort and the measurements will lead to both shirking and falsification of outcomes by the sensor. In this paper, we design an appropriate compensation scheme to incentivize the sensors to both exert costly effort and then reveal the resulting observation truthfully. To this end, we formulate the problem as a repeated game and propose a compensation scheme that employs stochastic verification by the operator coupled with an algorithm to assign a reputation to each sensor. By including the history of the behavior exerted by the sensor in determining present payments, we show that the operator can incentivize higher effort as well as more frequent truth-telling by the sensors.

Index Terms—Data acquisition, decision making, mechanism design, sensor network, state estimation.

I. INTRODUCTION

INCENTIVIZING individuals to follow policies desired by the system operator in smart sensor and infrastructure networks has received increased attention in recent years (see, e.g., [1]–[7] and the references therein). In such a setting, the system operator delegates tasks to several autonomous agents. Fulfilling these tasks may require costly effort by the agents, who may not benefit directly from the outcome of the task. Thus, these agents need to be incentivized to exert sufficient effort to complete tasks. Further, it may be costly, or even impossible, for the operator to assess the effort exerted by each agent and the accuracy of the reported outcomes; instead, she must rely on the data reported by the agent to do so. The focus of this paper is on designing a compensation scheme that incentivizes agents to both exert costly effort and reveal the outcomes truthfully in such settings.

Manuscript received September 3, 2018; revised May 2, 2019 and September 10, 2019; accepted September 26, 2019. Date of publication November 6, 2019; date of current version November 25, 2019. This work was supported in part by NSF under Grants ECCS-1446521 and CNS-1646019. The associate editor coordinating the review of this manuscript and approving it for publication was Prof. Laura Cottatellucci. (Corresponding author: Donya Ghavidel Dobakhshari.)

D. G. Dobakhshari and V. Gupta are with the Department of Electrical Engineering, University of Notre Dame, Notre Dame, IN 46556 USA (e-mail: dghavide@nd.edu; vgupta2@nd.edu).

P. Naghizadeh is with the Department of Electrical Engineering, Purdue University, West Lafayette, IN 47907 USA (e-mail: parinaz@purdue.edu).

M. Liu is with the Department of Electrical Engineering and Computer Science, University of Michigan, Ann Arbor, MI 48109 USA (e-mail: mingyan@umich.edu).

Digital Object Identifier 10.1109/TSP.2019.2952050

For concreteness, we focus on an estimation problem in which a central operator employs sensors to take measurements of a quantity of interest (e.g. a random variable or a parameter) at every time step over a predetermined time horizon, so as to generate an estimate of the quantity based on the sensors' reports. The sensors incur an effort cost for obtaining the measurements, which may model, e.g., the cost of operating the device, or battery power. The level of effort exerted by the sensors (and hence the true accuracy of their measurements) is their private information, and not known to the system operator. The operator aims to compensate sensors in a way that incentivizes them to exert high effort and truthfully report their measurements, so as to attain an accurate estimate of the quantity of interest.

Designing an appropriate contract for this setting is difficult due to two reasons: (i) profit misalignment in the sense that the sensors do not benefit directly from an accurate estimate at the operator, and (ii) information asymmetry between the operator and the sensor providing the information. To alleviate these challenges, the operator needs to design incentive mechanisms that mitigate both *moral hazard* (i.e., incentivizing desired actions by the sensors when effort is costly and private information for the sensor, see, e.g., [8, Chapter 4]) and *adverse selection* (i.e., incentivizing sensors to provide truthful information about the effort exerted when this information is private to them, see, e.g., [8, Chapter 3]).

While an extensive literature in contract theory (see, e.g., [1], [8], [9] and the references therein for an overview) has focused on resolving either moral hazard or adverse selection separately, the problem we consider features moral hazard followed by adverse selection in a *repeated setting*. This problem has received much less attention in the literature.

The closest works to ours are [10]–[13], which also consider the problem of moral hazard followed by adverse selection, although in *static* frameworks. In [10], the authors consider a binary information elicitation problem for multiple tasks when agents have endogenous proficiencies. The works of [11] and [12] explore the verification of outcomes generated by agents, directly and through comparison with reports of the other agents, respectively. Finally, [13] considers a setting featuring moral hazard followed by adverse selection for a single agent in a single stage interaction. A key difference of our work with [10]–[13] is in that these works have considered the interaction between the system operator and the agents in a *single* stage. The repeated setting that we consider herein adds additional dimensions to the problem, as it introduces the possibility of assigning reputations to the sensors based on the

history of their measurements, selecting payments based on both current reported measurements and sensors' reputations, and the selection of time-varying strategies by the sensors.

On the use of verification: We would first like to emphasize the need for intelligent use of verification in this crowdsensing problem. Given that the operator relies only on the information transmitted by the self-interested sensors, some form of verification or auditing of sensors' self-reported outcomes is crucial (see e.g., [14]–[16], in which verification is similarly proposed). Further, even if the operator can verify all sensors at all stages so as to ensure truthful revelation by the sensors through appropriate penalties, this approach may not, in general, be optimal for the operator due to the costly nature of verification.

We propose an alternative scheme by allowing verification to be done probabilistically at every time step. We then use this verification scheme, coupled with a *history-dependent* payment scheme, to design the contracts. In particular, we base the compensation on a *reputation score* assigned to each sensor. As such, the operator rewards the sensor based on the sequence of its actions (summarized by the reputation score) rather than merely on its behavior at the current stage. The operator assigns a higher reputation (and consequently higher payments) to a sensor that is verified and detected to be honest.

In terms of the methods through which verification or auditing can happen, we focus on *direct* verification, and interpret it as follows. In a sensor network application, we assume that the operator is able to deploy a trusted (but probably more costly) sensor in the field, or has the option of going directly to the physical location to collect her own measurements. This process of direct verification is indeed highly costly since it consumes time and effort from the operator. Such verification may further be imperfect or noisy. Our analysis in Section IV-D addresses this type of imperfect verification by the operator.

Similar notions of costly direct verification, known as “costly state verification,” appear in the economic literature on contract theory (see e.g. [17]). Specifically, costly state verification considers contract design problems in which verification (or disclosure) of enterprise performance is costly, and so a lender has to pay a monitoring cost to verify performance [18]. Our assumption of availability of direct costly verification is akin to this notion of costly state verification.

An alternative to direct verification in our context is cross-verification, where the operator verifies sensors' reports against each other. The use of cross-verification has been explored in crowdsourcing applications [19]–[23] as well as recently in sensor networks [24]. These works illustrate how the use of cross verification is in itself challenging, and at times ineffective, in soliciting truthful information; in particular, truthtelling, while an equilibrium, may be neither unique nor the maximum-reward equilibrium. Our study of direct verification allows us to forego these inefficiencies arising due to sensors' strategic behavior given cross-verification, and to focus solely on the effects of repeated interactions and reputation. The use of cross-verification in reputation-based contracts remains an interesting direction of future work.

Related literature: Using reputation for mitigating information asymmetry, particularly in repeated games, has been

commonly proposed, see e.g. [25]–[31]. In [25], the authors address the problem of mitigating pure adverse selection by means of reputation indices in a static setting. The authors of [26] present a comprehensive study of the use of reputation for mitigating adverse selection in repeated games. The works in [28]–[30] have focused on repeated interactions between a system operator and an agent under the assumption of only moral hazard for the agent. The work in [31] assumes knowledge of data quality as well as the sensors' cost. Specifically, it is assumed that the operator has a prior probability distribution on the quality of data, and that each sensor's cost is smaller than the (fixed) reward in a static setting. In contrast, we assume no prior knowledge on the quality and cost by the sensors, and further consider a repeated setup with both moral hazard and adverse selection.

The problem of providing incentives for participation in crowdsensing has also been considered in several papers. A comprehensive review of various incentive-based mechanisms, including both monetary and non-monetary incentives, is provided in [1]. Existing works have focused on incentives to mitigate moral hazard [24], [32], [33] or adverse selection [34], [35] that arise in crowdsensing. For non-verifiable outcomes, a class of peer prediction methods has been studied in [19]–[23]. These works use cross-verification between inputs of agents to make truthtelling a Nash equilibrium. However, these works assume knowledge of agent actions, and thus deal only with the problem of adverse selection. Additionally, truthtelling, while an equilibrium, is not necessarily the maximum-reward equilibrium in these mechanisms.

Our work is also related to the literature on optimal mechanism design. [36] reviews works on (the possibility of) optimal mechanism design for single stage interactions given private information to achieve properties like budget balance, individual rationality, and incentive compatibility. The work closest to ours on optimal mechanism design is [17], which studies the problem of optimal contract design using costly state verification when one party has private information. The setup of [17] considers (i) only private information and (ii) a single stage/static interaction between the two parties. In contrast, we consider both private information and private actions, as well as a repeated interaction setup.

Contributions: The main contribution of our work is to address the problem of simultaneously incentivizing high effort and truthful reports through contract design for repeated principal-agent interactions. For specificity, we focus on a distributed estimation setup as it may appear in a crowdsensing application; however, our techniques are applicable to any problem which exhibits moral hazard followed by adverse selection. We propose a reputation-based payment scheme coupled with stochastic verification for compensating the sensors. We show that under this scheme, compared to a payment scheme that does not use reputation scores, sensors will exert higher effort over time, and will truthfully disclose their accuracy with a higher frequency. Furthermore, the operator needs to resort to verification with a lower frequency. Nevertheless, the operator has to provide higher payments, as a result of which her overall payoff may decrease.

An earlier version of this work appeared in [37]. Compared to [37], this paper makes the following contributions: (i) our work in [37] studies the contract design problem for *two stage* interactions. Here, we extend our work to general multi-stage settings, and (ii) we consider several generalizations of the model to relax some of the assumptions made in [37]. In particular, we study the contract's dependence on the length of the interaction, analyze the effects of malicious sensors, consider an operator with a budget constraint, and evaluate the effects of imperfect verification by the operator.

Paper Organization: The remainder of the paper is organized as follows. In Section II, we present the model and some preliminaries. We analyze the proposed reputation-based payment scheme in Section III, and discuss its generalizations in Section IV. We conclude in Section V with some avenues for future work.

II. MODEL AND PRELIMINARIES

We study a repeated interaction between a system operator who is interested in estimating a quantity of interest, and I contracted agents (i.e., sensors) that can generate measurements at each time step over a finite horizon of length K .¹ We do not specify whether the quantity being reported is a random variable, parameter, or a field, since our setup is agnostic to that. The operator has no alternate sources for generating measurements. The accuracy of the measurements generated by the sensors increases with the effort they expend. Thus, the operator is interested in incentivizing high effort by the sensors, so as to attain sufficiently accurate estimates. However, both the true effort exerted by the sensors, as well as the outcomes they obtain in terms of the measurement or its accuracy, are unobservable by the operator. In other words, the operator faces moral hazard (in that she does not know the level of effort expended by the sensor) followed by adverse selection (in that she does not know the outcome of the effort). She therefore relies on the report by the sensors about the accuracy of their measurements.

Formally, the timeline is as follows. First, at time 0, the contract specifying the payment function between the operator and each sensor is signed. Then, at every stage/time k ($1 \leq k \leq K$), every sensor i ($1 \leq i \leq I$) performs the following actions:

- i) He exerts an effort $x_{ki} \in [0, \bar{x}_i]$ to generate a measurement. This effort incurs a cost of $g(x_{ki})$ for the sensor and leads to a measurement of accuracy level $\alpha(x_{ki})$.
- ii) He informs the operator of the measurements and the accuracy level that he generated. The sensor may misreport the accuracy level as the one corresponding to some other effort level $\hat{x}_{ki}$ as well as the measurements he generated.

After receiving measurements and accuracy reports from all the sensors at each time k , the operator verifies the reports from any sensor that she wishes to, by incurring a cost C per verification. For now, we assume that the verification is perfect, i.e., the operator can accurately detect any falsification by a sensor. We study the case of imperfect verification in Section IV-D. After

the verification, the operator makes payments as specified by the contract to each sensor.

We make the following assumption for simplicity.

Assumption 1: The accuracy $\alpha(x_{ki})$ is a deterministic function of x_{ki} , and the function is known to both the sensor and the operator. Therefore, without loss of generality, we may assume that each sensor reports simply his effort level to the operator.

The strategy space for each sensor is as follows. At each time step, sensor i decides on the pair $(x_{ki}, \hat{x}_{ki})$. If $\hat{x}_{ki} = x_{ki}$, we say that sensor i has been truthful (T) at stage k ; otherwise, we say that he has falsified the output and is non-truthful (NT). In general, the sensor will choose a mixed strategy; denote the probability that the sensor is truthful at stage k by q_{ki} .

On the other hand, the strategy for the operator consists of choosing the contract, which specifies the payment at every stage, and deciding at every stage whether or not to verify a given sensor. Denote the payment to sensor i at stage k by P_{ki} , which can depend on the reports from all the sensors till time k and the results of the verifications conducted by the operator till time k . The functional form of P_{ki} is specified in the contract. Denote the decision to verify a sensor (resp. not verify) by V (resp. by NV). Verification of sensor i at time k , captured through the variable $v_{ki} \in \{V, NV\}$, is done probabilistically; denote the probability that $v_{ki} = V$ by p_{ki} . Let z_{ki} denote the level of effort by sensor i known to the operator at the end of stage k . Since verification is perfect,

$$z_{ki} = \begin{cases} x_{ki} & \text{if } v_{ki} = V, \\ \hat{x}_{ki} & \text{if } v_{ki} = NV. \end{cases} \quad (1)$$

We shall denote the vector of efforts by the sensors at time k by $\mathbf{x}_k \triangleq \{x_{k1}, \dots, x_{kI}\}$, the set of all verification decisions at time k by $\mathbf{v}_k \triangleq \{v_{k1}, \dots, v_{kI}\}$ and that of all the payments at time k by $\mathbf{P}_k \triangleq \{P_{k1}, \dots, P_{kI}\}$.

We now specify the utility functions of the operator and the sensors. Sensors incur an effort cost in taking measurements and do not attach a value to the outcome of the task. Thus, the instantaneous utility of sensor i at stage k is given by,

$$U_{ki} = P_{ki} - g(x_{ki}). \quad (2)$$

On the other hand, the operator derives a benefit $S(x_k)$ from the measurements generated by the sensors at time k ; this depends on the (true) effort x_{ki} of the sensors. We assume that this function is increasing and concave. Without loss of generality, we normalize $S(0) = 0$, i.e., the operator does not derive any benefit when $x_{ki} = 0, \forall i, k$. The costs that the operator incurs are from the payments she pays to the sensors and from any verification that she performs which is denoted by C . Thus, the instantaneous payoff (or utility) of the operator at stage k is given by,

$$\Pi_k = S(\mathbf{x}_k) - \sum_{i=1}^I (P_{ki} + C \mathbb{1}\{v_{ki} = V\}), \quad (3)$$

where $\mathbb{1}\{v_{ki} = V\} = \begin{cases} 1 & \text{if } v_{ki} = V \\ 0 & \text{otherwise} \end{cases}$. We assume that both the sensors and the operator discount future payoffs with a factor δ ,

¹We will henceforth use she/her for the operator, and he/his for the sensor.

so that their payoffs over the entire time horizon are given by,

$$U_i \triangleq \sum_{k=1}^K \delta^k U_{ki}, \Pi \triangleq \sum_{k=1}^K \delta^k \Pi_k. \quad (4)$$

In this paper, we consider the design of the payment vectors $P_1, P_2, \dots, P_K$ and the verification decisions $v_1, v_2, \dots, v_K$ by the operator so that rational self-interested sensors will take actions that lead to maximization of the operator's utility Π . We impose two constraints on the problem.

- i) Incentive compatibility (IC): IC is a standard constraint imposed in a contract to align the effort by the sensor to the effort desired by the operator. A contract is incentive compatible if every sensor chooses to take the action preferred by the operator. Note that the sensor is a rational decision maker, and therefore chooses his effort level to maximize his expected payoff $\mathbb{E}[U_i]$. Thus, we impose that if x_{ki}^* denotes the effort desired by the operator at time k by agent i , then the compensation scheme should be such that $\mathbb{E}[U_i(x_{ki}^*)] \geq \mathbb{E}[U_i(x_{ki})], \forall x_{ki}$. Note that the expectation is with respect to the verification strategy of the operator.
- ii) Individual rationality (IR) or participation constraints: Both the operator and the sensors should prefer participation in the proposed scheme to opting out. In other words, their expected continuation utility at each time t (defined as the expected utility of the sensors and the operator at each time t looking forward till time K) must be greater than the utility achieved by opting out.² Formally, assuming that the operator requires the participation by all the sensors, we impose

$$\mathbb{E} \left[\sum_{k=t}^K \delta^{k-t} \Pi_k \right] \geq 0, \quad \mathbb{E} \left[\sum_{k=t}^K \delta^{k-t} U_{ki} \right] \geq 0, \quad \forall t, i.$$

Note that the expectation in the first term (the continuation utility of the operator) is with respect to the strategy of reporting truthfully or not by the sensor.

The optimization problem for the operator that we seek to solve is therefore given by

$$\mathcal{P} : \begin{cases} \max_{\{v_1, \dots, v_K\}, \{P_1, \dots, P_K\}} \mathbb{E}[\Pi] \\ \text{s.t. IC and IR constraints.} \end{cases}$$

Problem $\mathcal{P}$ is, in general, difficult to solve given the freedom in choosing the payment function as a function of the entire history of the actions of the sensors and the operator. Instead, we simplify the problem by assuming that the history of the past reported efforts of the sensors and the results of the verification done by the operator are summarized through a *reputation score*, that is then used to design the payment scheme.

Specifically, assume that the operator assigns a *reputation score* R_{ki} to sensor i at beginning of stage k based on the history of the sensor's past reputation scores, as well as the reported

effort z_{ki} as known to the operator at stage k , i.e.,

$$R_{ki} = f_k(R_{1i}, \dots, R_{(k-1)i}, z_{ki}), k = 1, \dots, K, \quad (5)$$

where $R_{0i} = 0$ and the function $f_k(\cdot)$ is the *reputation function* selected by the operator at stage k . The operator then uses these reputation scores to offer monetary compensations $P_{ki} := P(R_{ki})$ to the sensors based on a payment function $P(\cdot)$ that is a design choice of the operator. In this paper, we focus on the class of *weighted reputation* of the following format:

$$\begin{cases} R_{1i} = R(z_{1i}), \\ R_{ki} = \omega R(z_{ki}) + (1 - \omega)R_{(k-1)i}, k = 2, \dots, K \end{cases} \quad (6)$$

where by an abuse of notation, we once again use $R(\cdot)$ to denote the (instantaneous) reputation function, $\omega \in [0, 1]$ is the *reputation weight*, and z_{ki} is given by (1).

A few comments about this reputation function are in order:

- By adjusting the value of ω , the operator can place varying importance on the history of past behavior of the sensor in assessing the current payment. For instance, when $\omega = 1$, the compensation is based solely on the (reported or verified) effort expended at the current stage. We refer to this special case of compensation scheme as *instant payments*.
- By the definition of z_{ki} in (1), if the sensor is not verified, $R(\cdot)$ is evaluated based on the output $\hat{x}_{ki}$ reported by the sensor at that stage. If the sensor is verified, he is assigned a reputation based on his verified output x_{ki} .
- The form of the reputation function implies that there is no cross-verification done among sensors. Cross-verification increases the problem complexity significantly and is left for future work.

Further, $R(\cdot)$ is assumed to be concave. Further, we let the minimum and maximum of reputation function be $R(0) = l$ and $R(\bar{x}) = h$, where $\{l, h\} \geq 0$ are designer specified parameters.

We propose the payment at each stage k for sensor i to be equal to the reputation of the sensor at that stage, i.e., $P_{ki} = R_{ki}$. Thus, the instantaneous utility of sensor i at stage k is now given by

$$U_{ki} = R_{ki} - g(x_{ki}), \quad (7)$$

while that of the operator is

$$\Pi_k = S(x_k) - \sum_{i=1}^I (R_{ki} + C \mathbb{1}\{v_{ki} = V\}). \quad (8)$$

As before, participants discount the future by a factor δ , and their long-run utilities U_i and Π can be defined as in (4).

Remark 1: Note that the use of verification is indispensable: if a sensor is not verified, he will always exert effort $x_{ki} = 0$, and falsify his effort as $\hat{x}_{ki} = \bar{x}_i$. Nevertheless, since verification is costly, the goal of introducing a reputation-based payment scheme is to reduce the verification frequency.

III. MAIN RESULTS

We now consider the K -stage game between the sensors and the operator by solving problem $\mathcal{P}$ with payments given by (6). In this section, we proceed under the following assumptions; we will relax Assumption 2 in Section IV-C.

²Without loss of generality we assume that the utilities of the sensors and the operator when they opt out are zero.

TABLE I
INSTANTANEOUS PAYOFFS OF THE SENSOR (ROW PLAYER) AND THE OPERATOR
(COLUMN PLAYER) AT EACH STAGE k

	V	NV
T	$R(x_k) - g(x_k), S(x_k) - R(x_k) - C$	$R(x_k) - g(x_k), S(x_k) - R(x_k)$
NT	$R(0), S_k^0 - R(0) - C$	$R(\bar{x}), S_k^0 - R(\bar{x})$

Assumption 2: Throughout this section, we assume no budget constraint for the operator, i.e., $\sum_{i=1}^I P_{ki}$ can be chosen arbitrarily.

Assumption 3: We assume that sensors inputs are equivalent from the operator's perspective in the sense that $S(\cdot)$ is independent of the permutation of sensors' indices for zero and maximum effort. Formally, let S_{ki}^0 be the operator's benefit at stage k when sensor i exerts zero effort $S_{ki}^0 \triangleq S(x_{k1}, \dots, x_{k(i-1)}, 0, x_{k(i+1)}, \dots, x_{kI})$. Similarly, define $\bar{S}_{ki} \triangleq S(x_{k1}, \dots, x_{k(i-1)}, \bar{x}, x_{k(i+1)}, \dots, x_{kI})$. We assume that value of S_{ki}^0 and $\bar{S}_{ki}$ are independent of i and constant for a given k , i.e., $S_{ki}^0 = S_k^0$ and $\bar{S}_{ki} = \bar{S}_k$, for any i .

Under Assumptions 2 and 3 and in the absence of cross-verification, the payments to various sensors as well as the benefit of the operator from measurements are effectively decoupled. Thus, we first restrict attention to studying the contract design problem for a single sensor without loss of generality. For notational ease, we drop the subscripts i .

A. The Payoff Matrix

To find the payoffs of the sensor and the operator, note that we have assumed no falsification cost for the sensor. We have also assumed that the reduction in reputation due to any detected falsification is independent of the amount of falsification. As a result, for the stage game, if the sensor behaves strategically at stage k , he does not exert any effort and realizes $x_k = 0$, but reports the maximum effort $\hat{x}_k = \bar{x}$, to gain the maximum reputation/payment if not verified. The payoff matrix of the stage game is thus specified in Table I. For the game in Table I, the operator designs the parameters of the compensation scheme, to satisfy the IC and IR constraints, and maximize her profit.

We now proceed to the analysis of the K -stage game with stage games specified in Table I.

B. Nash Equilibria and Optimal Choice of Effort

We start with the pure strategy Nash Equilibrium (NE) of the game.

Proposition 1: The only pure strategy Nash equilibrium of the K -stage game with stage games specified in Table I is repeated play of the strategy (NT, NV).

Proof: See Appendix. ■

The pure strategy equilibrium identified in Proposition 1 is such that the operator offers no payment to the sensor, and the sensor exerts no effort. This is equivalent to the outside option for the operator. We, therefore, consider the mixed strategy equilibria of the game. The following theorem characterizes the mixed strategy NE of the game in Table I.

Theorem 1: Under the weighted reputation scheme in (6),

- i) the mixed strategy equilibrium of the game at each stage k (if exist) in Table I is unique and as follows. The operator verifies the sensor with probability

$$p_k = \frac{h - R(x_k) + \frac{g(x_k)}{\Gamma(k)}}{h - l},$$

and the sensor reveals the truth at stage k with probability

$$q_k = \frac{h - l - \frac{C}{\Gamma(k)}}{h - l},$$

$$\text{where } \Gamma(k) = \begin{cases} \sum_{j=k}^K \delta^{j-k} (1 - \omega)^{j-k} & k = 1 \\ \omega \sum_{j=k}^K \delta^{j-k} (1 - \omega)^{j-k} & k \neq 1 \end{cases}$$

- ii) The sensor chooses $\tilde{x}_k$ such that

$$\Gamma(\omega) \frac{\partial R(x_k)}{\partial x_k} = \frac{\partial g(x_k)}{\partial x_k} \Big|_{x_k = \tilde{x}_k}.$$

Proof: See Appendix. ■

Remark 2: For these mixed strategy equilibria to exist, the operator should select h , ω , and l such that, at each k ,

$$\Gamma(k) \geq \max \left\{ \frac{g(x_k)}{R(x_k) - l}, \frac{C}{h - l} \right\}. \quad (9)$$

The operator has to take this condition into account in the contract design. Otherwise, the only possible outcome will be the pure strategy Nash equilibrium (NT, NV).

Note that $\Gamma(k+1) < \Gamma(k)$. As a result, when the game approaches the latter stages, the sensor chooses lower q_i , i.e., the sensor is truthful with a lower probability. The operator on the other hand, even for the same effort level x_k , verifies the sensor with higher probability p_i as the game progresses. This can be intuitively interpreted as the fact that maintaining high reputation scores become less attractive near the contract's terminal stages. Similarly, for a fixed k , as K increases, Γ increases. As a result, falsification and verification probabilities are higher at a fixed time with shorter interactions (smaller K).

C. Optimal Choice of Payment Parameters

In order to further solve for the optimal choices of the contract parameters, we proceed with the analysis under the following two assumptions. First, we assume a linear model on both the cost function of the sensor and the reputation function.

Assumption 4: The reputation and cost functions are linear. Specifically,

- the reputation function is given by $R(x) = \frac{h-l}{\bar{x}}x + l$, and
- the cost of effort is $g(x_k) = bx_k$.

The second assumption introduces an upper bound on the value of $\bar{S}_k$, the attainable benefit of the operator from the sensor maximum effort.

Assumption 5: We assume that $0 < C(\bar{S}_k - S_k^0) < (b\bar{x})^2$ for each stage k .

Remark 3: Note that the above inequality can be re-written as $\frac{b\bar{x}}{\bar{S}_k - S_k^0} > \frac{C}{b\bar{x}}$. For a normalized $C = 1$, we obtain, $b\bar{x} > \sqrt{(\bar{S}_k - S_k^0)}$. Assumption 5 therefore imposes that the cost of effort to realize the maximum accuracy is more than benefit of the operator due to that accuracy.

Note that under Assumption 4, the operator chooses h , l , and ω to maximize $\mathbb{E}[\Pi]$, subject to the (IC) and (IR) constraints of the sensor as stated in Problem $\mathcal{P}$. We start by optimizing the choice of l and h given a fixed reputation weight ω .

Theorem 2: Under Assumption 4 and for a given ω ,

- i) the individual rationality constraint for the sensor is always satisfied at each stage k , and
- ii) the optimal value of the lowest reputation score is $l^* = 0$.

Proof: See Appendix. ■

Theorem 3: Under Assumptions 4 and 5,

- i) the optimal value of h is given by $h^* = \frac{b\bar{x}}{\omega}$, and
- ii) for $\omega \neq 1$, the operator will incentivize effort level $\bar{x}$ at each stage $k \neq K$, and the effort x^* at the last stage, where $\frac{\partial S(x)}{\partial x}|_{x=x^*} = b$.
- iii) The verification probabilities for the operator p_k and the probabilities q_k for the sensor to be truthful are given by

$$p_k = \frac{\bar{x} + x_k \left(\frac{\Gamma(K)}{\Gamma(k)} - 1 \right)}{\bar{x}}, \quad q_k = \frac{b\bar{x} - \frac{C}{\Gamma(k)}}{b\bar{x}}.$$

Proof: See Appendix. ■

Next, we study the role of reputation weight i.e., importance of history in the proposed contract. Intuitively, the reputation weight ω determines the importance of inter-temporal incentives (i.e., conditioning future payments on the history of past efforts). In particular, $\omega = 1$ yields an instant payment scheme, in which no inter-temporal incentives are present. For this case, the actions of the operator and the sensor are as follows.

Corollary 1: If $\omega = 1$, the sensor realizes the effort profile $(x^*, \dots, x^*)$ where $\frac{\partial S(x)}{\partial x}|_{x=x^*} = b$, the operator verifies the sensor with probability $p = 1$ at each stage, and the sensor is only truthful with probability $q = 1 - \frac{C}{b\bar{x}} < 1$.

Proof: The proof is similar to Theorem 3 for $\omega = 1$ and is omitted. ■

By comparing Theorem 3 and Corollary 1, we observe that while the verification frequency, falsification probabilities, and the effort level of the sensor, with and without the use of reputation, remain the same at the last stage, the values at the other stages differ due to the introduction of inter-temporal incentives. In particular, when history-dependent reputations are used, the operator needs to verify the sensor with a lower probability, and the sensor is truthful with a higher probability. Furthermore, the sensor exerts higher effort in the non-terminal stages.

We next consider the optimal choice of ω , under which expected payoff of the operator is maximized.

Theorem 4: Under Assumption 4 and 5, a choice of $\omega = 1$ maximizes the operator's payoff. That is, instant payments yield higher payoffs than payments based on linearly weighted reputations.

Proof: See Appendix. ■

We observe that while using inter-temporal incentives through linearly weighted reputation functions can benefit the operator by reducing the required verification frequency, increasing the effort level of the sensor, and increasing the probability of truthfulness, it will nevertheless reduce the overall payoff of the operator. This is because the operator has to now offer a higher compensation to the sensor under reputation-based payments.

TABLE II

PAYOFFS OF THE MALICIOUS SENSOR (ROW PLAYER) AND THE OPERATOR (COLUMN PLAYER) AT EACH STAGE k FOR THE GAME IN SECTION IV-B

	V	NV
T	$R(x_k) - bx_k, S(x_k) - R(x_k) - C$	$R(x_k) - bx_k, S(x_k) - R(x_k)$
NT	$R(0), -R(0) - C$	$R(\bar{x}) + \alpha, -R(\bar{x}) - \alpha$

IV. EXTENSIONS

In this section, we consider several generalizations of the model. In particular, we study the behavior of the sensors when the number of stages grows, when a sensor is malicious, when the operator has a budget constraint, and when verification by the operator is imperfect. For ease of notation, we will set $l = 0$ and $S_k^0 = 0$ for this section and in subsection IV-B, IV-C, and IV-D, we assume $K = 2$.

A. The Infinite-Stage Game

We first consider the role of inter-temporal incentives in an infinitely repeated interaction.

Corollary 2: In the infinite-stage game as $K \rightarrow \infty$,

- i) the probabilities of verification p_k and truthtelling q_k are given by

$$p_k = \frac{h - R(x_k) + \frac{g(x_k)}{\Gamma}}{h}, \quad q_k = \frac{h - \frac{1}{\Gamma}C}{h}, \quad (10)$$

where $\Gamma = \frac{\omega}{1 - (1 - \omega)\delta}$.

- ii) At each stage, the sensor chooses $\tilde{x}$ such that

$$\Gamma \frac{\partial R(x_k)}{\partial x_k} = \frac{\partial g(x_k)}{\partial x_k} \Big|_{x_k = \tilde{x}}.$$

- iii) Under Assumption 4, the sensor chooses $\bar{x}$ at each stage,
- iv) under Assumption 5, the optimal value of h is $h = \frac{b\bar{x}}{\Gamma}$ and the operator is indifferent among all choices of ω .

Proof: See Appendix. ■

The above proposition shows that as $K \rightarrow \infty$, the sensor behaves consistently and puts in the maximum effort at each stage. Further, the utility of the operator is independent of the choice of ω . In other words, the operator tends to be indifferent about employing reputation-based payments.

B. Malicious Sensors

In this section, we analyze the behavior of a malicious sensor under the proposed contract scheme. For simplicity, we focus on the two-stage problem with a linear cost function for the sensor. We define a malicious sensor as a sensor who gains an extra benefit, denoted by α , when he falsifies the transmitted data and this is gone undetected. In particular, when untruthful, this malicious sensor falsifies both the reported measurement and its accuracy, claiming maximum accuracy for the falsified measurement. By doing both types of falsification, the sensor can maximally misguide the estimation at the operator. The payoff matrix of the stage game is thus specified in Table II.

We observe the following behavior by the sensor and the operator at the mixed strategy Nash equilibrium of this game.

Proposition 2: Under the weighted reputation scheme, the mixed strategy equilibrium of the game in Table II is as follows.

The operator verifies the sensor with probabilities

$$p_2 = \frac{h + \alpha - R(x_2) + \frac{g(x_2)}{\omega}}{h + \alpha},$$

$$p_1 = \frac{h + \alpha - R(x_1) + \frac{g(x_1)}{1+(1-\omega)\delta}}{h + \alpha}.$$

and the sensor reveals the truth with probabilities

$$q_2 = \frac{h + \alpha - \frac{1}{\omega}C}{h + \alpha}, \quad q_1 = \frac{h + \alpha - \frac{1}{1+(1-\omega)\delta}C}{h + \alpha}.$$

If $b\bar{x} > \sqrt{C(S(\bar{X}) - S_{0_i} + \delta^\omega \alpha)}$ and Assumption 4 holds

- i) the optimal reputation parameter is given by $h = \frac{b\bar{x}}{\omega} - \alpha$,
- ii) the operator will incentivize effort level $\bar{x}$ at the first stage, and the effort 0 at the second stage, and
- iii) the optimal value of ω is given by the solution to the following

$$-b\bar{x} \frac{\partial f(\omega)}{\partial \omega} \left(1 - \frac{1}{f^2(\omega)} \frac{S(\bar{x})C}{(b\bar{x})^2} \right) - \alpha \left(\delta + \frac{C}{b\bar{x}} \right) = 0,$$

where $f(\omega) := \frac{\delta^\omega}{\omega}$.

Proof: The proof is similar to those for contract design with a strategic sensor and is omitted. ■

The above proposition shows that in the presence of a malicious sensor, both the verification probabilities and the truth-telling probabilities increase compared to the case with a non-malicious sensor. While the former is to be expected, the latter may be non-intuitive. However, one interpretation is as follows: by maintaining a high enough probability of truth-telling, the malicious sensor can reduce the verification probability, and instead gain α whenever verification is not invoked.

Further, we can observe that at equilibrium in the first stage, both malicious and strategic sensors behave similarly in terms of the actions they take; however, the malicious sensor takes zero action in the second stage. Note that to incentivize effort level $\bar{x}$ at the first stage, the operator has to verify the malicious sensor with a higher probability in both stages. In other words, the operator will incur higher monitoring costs when dealing with a malicious sensor; nevertheless, the added expenditure allows her to recover the same accuracy of measurements as with a non-malicious sensor.

C. Budget Constraint at the Operator

Consider two sensors participating in the crowdsensing mechanism, and suppose that the operator has a limited budget to compensate the sensors.³ In this case, in addition to seeking compensation by providing accurate measurements, the sensors are in competition with each other. Thus, we consider the compensation of sensor i to be a function of not only his own report, but also the accuracy reported by the other sensor. Denote the payment to the first and second sensors at stage k by $R_1(z_{k1}, z_{k2})$ and $R_2(z_{k1}, z_{k2})$, respectively. The payoff matrix of the stage game is specified in Table III.

³The generalization to the case of I sensors is straightforward.

TABLE III
PAYOFFS OF THE SENSOR i (ROW PLAYER) AND THE OPERATOR (COLUMN PLAYER) AT EACH STAGE k FOR THE GAME IN SECTION IV-C

	V	NV
T	$R_i(\mathbf{x}_k) - b x_{ki}, S(\mathbf{x}_k) - R_i(\mathbf{x}_k) - C$	$R_i(X_k) - b x_{ki}, S(\mathbf{x}_k) - R_i(\mathbf{x}_k)$
NT	$R_i(0, x_{k(-i)}), -R_i(0, x_{k(-i)}) - C$	$R_i(\bar{x}, x_{k(-i)}), -R_i(\bar{x}, x_{k(-i)})$

We assume that the operator's budget is given by c , so that she is constrained to $R_1(z_{k1}, z_{k2}) + R_2(z_{k1}, z_{k2}) = c$. Further, following the classical Cournot game setup, we propose the payments of the following form to each sensor i at stage k

$$R_i(z_{k1}, z_{k2}) = z_{ki} \left(\lambda - \sum_{j=1}^2 z_{kj} \right). \quad (11)$$

The proposed payment/reputation captures the budget constraint at the operator. To see this, note that the payment is increasing in the separated effort for small value of the efforts, but decreasing for larger values. The operator is therefore discouraging large collective effort, as she faces a budget constraint and cannot compensate the sensors if both exert high effort. Note also that the proposed payment indicates that when the overall level of accuracy reported by sensors exceeds the budget threshold λ , the sensors will be penalized by the operator.⁴ The reputation function is assumed to be given by

$$\begin{cases} R_{1i} = R_i(z_{11}, z_{12}), \\ R_{2i} = \omega R_i(z_{21}, z_{22}) + (1 - \omega) R_{1i} \end{cases}, \quad (12)$$

where the reputation weight $\omega \in [0, 1]$.

Proposition 3: Under the weighted reputation scheme (12) and linear cost functions for the sensors, the unique mixed strategy equilibria of the game in Table III are as follows.

- i) The operator verifies sensor i with probabilities⁵

$$p_{2i} = \frac{R_i(\bar{x}, x_{k(-i)}) - R_i(x_k) + \frac{1}{\omega} b x_2}{R_i(\bar{x}, x_{k(-i)})},$$

$$p_{1i} = \frac{R_i(\bar{x}, x_{k(-i)}) - R_i(x_k) + \frac{1}{\delta^\omega} b x_1}{R_i(\bar{x}, x_{k(-i)})}.$$

and the sensors reveal the truth with probabilities

$$q_{2i} = \frac{R_i(\bar{x}, x_{k(-i)}) - \frac{1}{\omega} C}{R_i(\bar{x}, x_{k(-i)})}, \quad q_{1i} = \frac{R_i(\bar{x}, x_{k(-i)}) - \frac{1}{\delta^\omega} C}{R_i(\bar{x}, x_{k(-i)})},$$

where $\delta^\omega = 1 + (1 - \omega)\delta$ and $\mathbf{x}_k = (x_{k1}, x_{k2})$.

- ii) At each stage, the optimal effort level chosen by the sensors are equal, and are given by $x_{1i} = \frac{\lambda - \frac{b}{\delta^\omega}}{3}$ and $x_{2i} = \frac{\lambda - \frac{b}{3}}{3}$ respectively.
- iii) The optimal value of ω is $\omega^* = 1$.

Proof: The proof follows steps similar to that of Theorem 1 and 4, with Table III instead of Table I as the stage game and is omitted. ■

⁴Note that $R_1(z_{k1}, z_{k2}) + R_2(z_{k1}, z_{k2}) = (\sum_{j=1}^2 z_{kj})(\lambda - \sum_{j=1}^2 z_{kj})$, with the maximum happening at $\sum_{j=1}^2 z_{kj} = \frac{\lambda}{2}$ which yields $\frac{\lambda^2}{4} = c$.

⁵Here we have borrowed notation from game theory literature and used $x_{k(-i)}$ to represent the effort by every sensor at stage k except for sensor i .

TABLE IV
PAYOFFS OF THE SENSOR (ROW PLAYER) AND THE OPERATOR (COLUMN
PLAYER) AT STAGE k FOR IMPERFECT VERIFICATION WITH PROBABILITY μ FOR
THE GAME IN SECTION IV-D

	V	NV
T	$R_k(x_k) - g(x_k), S(x_k) - R_k(x_k) - C$	$R_k(x_k) - g(x_k), S(x_k) - R_k(x_k)$
NT	$\mu R_k(0) + (1 - \mu)R_k(\bar{x}), -\mu R_k(0) - (1 - \mu)R_k(\bar{x}) - C$	$R_k(\bar{x}), -R_k(\bar{x})$

We observe that similar to the case without a budget constraint, the optimal effort in the last stage is lower compared to the previous stage. In addition, we observe that due to the budget constraint, the sensors end up exerting lower effort compared to the case without budget constraints. Calculating the expected utility of the operator as a function of λ , we also observe that as the value of λ increases, i.e., as we relax the budget constraint, the expected utility of the operator decreases. In other words, as the budget constraint for the operator becomes less constraining, the compensation to the sensors increases and the expected utility of the operator declines.

D. Imperfect Verification by the Operator

So far we have assumed the verification of efforts of the sensors by the operator is perfect, in the sense that the operator can accurately detect any falsification by a sensor. In this section, we analyze the case when this verification is not perfect. In particular, consider a set-up in which the verification done by the operator has accuracy μ . That is, there is a probability $1 - \mu$ that an operator cannot detect falsification during verification. Under this definition, the game in Table I will be modified to Table IV; note that the utility of the operator and the sensor changes only for (NT, V) , while all other utilities remain the same. In the following proposition, we present the Nash Equilibrium attained under imperfect verification.

Proposition 4: When imperfect verification is implemented by the operator, mixed strategy equilibrium of the game in Table IV at each time is unique and as follows. The operator verifies the sensor with probabilities

$$p_2 = \frac{h - R(x_2) + \frac{bx_2}{\omega}}{\mu h}, \quad p_1 = \frac{h - R(x_1) + \frac{bx_1}{1 + (1 - \omega)\delta}}{\mu h}.$$

and the sensor reveals the truth with probabilities

$$q_2 = 1 - \frac{C}{\mu h \omega}, \quad q_1 = 1 - \frac{C}{\mu h (1 + (1 - \omega)\delta)}.$$

Proof: The proof follows steps similar to that of Theorem 1, with Table IV instead of Table I as the stage game and is omitted. ■

We can observe that compared to the case with perfect verification, imperfect verification leads to higher probability of verification and lower probability of truthful behavior by the sensor. Further, calculating utilities of the sensor and operator, we observe that while the utility of the sensor remains the same, the utility of the operator is lower under imperfect verification. ■

V. CONCLUSION

In this paper, we studied the problem of contract design between a system operator and strategic sensors in a repeated setting. The sensors are asked to exert costly effort to collect sufficiently accurate observations for the operator. As the effort invested and the accuracy of the resulting outcome are both private information of the sensors, the operator needs to design a compensation scheme that mitigates moral hazard followed by adverse selection. We proposed a reputation-based payment scheme coupled with stochastic verification. We showed that by increasing the importance of past behavior in our proposed linearly weighted reputation-based payments, the sensors exert higher effort, and have a higher probability of being truthful. The operator, on the other hand, can invoke verification less frequently. However, the operator offers higher payments to the sensors, which leads to a lower overall profit.

We have so far considered inter-temporal incentives that are based on a linearly weighted reputation function. Considering other functional forms for evaluating reputations, and its impact on the operator and sensors strategies and payoffs, is an important direction of future work. In addition, we have considered the design of individual contracts for each sensor, due to our assumption of independent measurements by the sensors. As an interesting direction of future work, we are interested in analyzing the contract design problem for multiple sensors when the outcomes of the estimate at the sensors are coupled. Such coupling may enable the operator to further cross-verify the outcomes of the sensors as part of the payment scheme.

APPENDIX

Proof of Proposition 1: We start with the potential pure Nash equilibrium (T, V) , and analyze the payoff of the operator. By playing V over NV at the first stage, the operator decreases her utility by C at the subsequent stage. Therefore, (T, NV) dominates (T, V) . Similarly, by analyzing the payoff of the sensor, we can see that (NT, NV) dominates (T, NV) . We have therefore discarded (T, V) and (T, NV) as pure Nash equilibria of the stage games. Therefore, if a pure strategy NE exists, the sensor will be playing NT . However, given that the sensor is always playing NT , the operator's optimal choice is to set $h = 0$, and play NV . Therefore, the only possible pure strategy NE of the game is (NT, NV) , i.e., the operator offers no payment to the sensor which is in fact the operator's outside option. ■

Proof of Theorem 1: (i) We use backward induction to find the operator and sensor's strategies. Further, denote by U_k^c the expected continuation utility of the sensor from time k looking forward, i.e.,

$$U_k^c = \mathbb{E} \left[\sum_{j=k}^K \delta^{j-k} U_j \right].$$

Prior to presenting the mixed strategy equilibrium at each stage k , we first prove the following lemma on the utility of the sensor. ■

Lemma 1: The expected continuation payoff of the sensor at stage k in the mixed strategy equilibrium is given by

$$U_k^c = \sum_{j=k}^K \delta^{j-k} \left((1-\omega)^{j-k+1} R_{k-1} + \Gamma(j)R(x_j) - g(x_j) \right). \quad (13)$$

Further, the total expected utility of the sensor over the entire K stages is given by

$$\mathbb{E}[U] = \sum_{k=1}^K \delta^{k-1} (\Gamma(k)R(x_k) - g(x_k)). \quad (14)$$

where

$$\Gamma(k) = \begin{cases} \omega \sum_{j=k}^K \delta^{j-k} (1-\omega)^{j-k} & k \neq 1 \\ \sum_{j=k}^K \delta^{j-k} (1-\omega)^{j-k} & k = 1 \end{cases}.$$

Proof: We use backward induction to prove that the expected continuation payoff of the sensor follows (13). First note that since we assume linear payment functions based on the sensor's reputation, the (expected continuation) utility of the sensor at stage k is given by

$$U_k^c = (1-\omega)R_{k-1} + \omega R(x_k) - g(x_k) + \delta U_{k+1}^c. \quad (15)$$

The sensor's utility at the last stage stage K is then as follows

$$U_K^c = (1-\omega)R_{K-1} + \omega R(x_K) - g(x_K).$$

The utility of the sensor at stage $K-1$ is therefore given by,

$$\begin{aligned} U_{K-1}^c &= (1-\omega)R_{K-2} + \omega R(x_{K-1}) - g(x_{K-1}) + \delta U_K^c \\ &= (1-\omega)R_{K-2} + \omega R(x_{K-1}) - g(x_{K-1}) + \delta \omega R(x_K) \\ &\quad - \delta g(x_K) + \delta(1-\omega)R_{K-1}. \end{aligned}$$

Given (6) on the reputation score update, note that

$$R_{K-1} = \omega R(x_{K-1}) + (1-\omega)R_{K-2}.$$

Thus, U_{K-1}^c can be written as

$$\begin{aligned} U_{K-1}^c &= ((1-\omega) + \delta(1-\omega)^2)R_{K-2} \\ &\quad + (\omega + \omega\delta(1-\omega))R(x_{K-1}) - g(x_{K-1}) \\ &\quad + \delta(\omega R(x_K) - g(x_K)), \end{aligned}$$

which yields

$$\begin{aligned} U_{K-1}^c &= \sum_{j=K-1}^K \delta^{j-K+1} \\ &\quad \times \left((1-\omega)^{j-K+2} R_{K-2} + \Gamma(j)R(x_j) - g(x_j) \right). \end{aligned}$$

which follows (13), forming the induction basis. Next, we assume that U_k^c follows (13) and then prove U_{k-1}^c also follows (13). Assume the utility of the sensor at stage k is given by

$$U_k^c = \sum_{j=k}^K \delta^{j-k} \left((1-\omega)^{j-k+1} R_{k-1} + \Gamma(j)R(x_j) - g(x_j) \right).$$

Based on (15), U_{k-1}^c is derived as

$$U_{k-1}^c = (1-\omega)R_{k-2} + \omega R(x_{k-1}) - g(x_{k-1}) + \delta U_k^c.$$

Replacing U_k^c with its expansion yields,

$$\begin{aligned} U_{k-1}^c &= (1-\omega)R_{k-2} + \omega R(x_{k-1}) - g(x_{k-1}) \\ &\quad + \sum_{j=k}^K \delta^{j-k+1} \left((1-\omega)^{j-k+1} R_{k-1} \right. \\ &\quad \left. + \Gamma(j)R(x_j) - g(x_j) \right). \end{aligned}$$

Given $R_{k-1} = \omega R(x_{k-1}) + (1-\omega)R_{k-2}$, we can write

$$\begin{aligned} U_{k-1}^c &= (1-\omega)R_{k-2} + \omega R(x_{k-1}) - g(x_{k-1}) \\ &\quad + \sum_{j=k}^K (1-\omega)^{j-k+1} \delta^{j-k+1} \left(\omega R(x_{k-1}) + (1-\omega)R_{k-2} \right) \\ &\quad + \sum_{j=k}^K \delta^{j-k+1} (\Gamma(j)R(x_j) - g(x_j)). \end{aligned} \quad (16)$$

Note that

$$\begin{aligned} &\omega R(x_{k-1}) + \sum_{j=k}^K (1-\omega)^{j-k+1} \delta^{j-k+1} \omega R(x_{k-1}) \\ &= \sum_{j=k-1}^K (1-\omega)^{j-k+1} \delta^{j-k+1} \omega R(x_{k-1}) = \Gamma(k-1)R(x_{k-1}). \end{aligned}$$

Thus, if we update (16), U_{k-1}^c is given by

$$U_{k-1}^c = \sum_{j=k-1}^K \delta^{j-k+1} \left((1-\omega)^{j-k+2} R_{k-2} + \Gamma(j)R(x_j) - g(x_j) \right),$$

which is consistent with (13). This completes the proof of the first part of the Lemma 1. For the second part, note that the long-run utility of the sensor over the entire horizon is calculated by substituting k with 1, and noting that $R_0 = 0$, leading to,

$$\mathbb{E}[U] = \sum_{k=1}^K \delta^{k-1} (R(x_k)\Gamma(k) - g(x_k)).$$

■

Assume that the operator verifies the sensor with probability p_k . If the sensor reports truthfully, i.e., chooses T , at stage k , his expected utility is given by

$$U_k^c(R_{k-1}, R(x_k), \dots, R(x_K)). \quad (17)$$

We next consider the payoff from falsification by the sensor, i.e., playing NT . Recall that when the sensor plays NT , he also exerts no effort, i.e., $x_k = 0$. Therefore, the expected utility from non-truthful behavior is given by

$$p_k U_k^c(R_{k-1}, l, \dots, R(x_K)) + (1-p_k) U_k^c(R_{k-1}, h, \dots, R(x_K)). \quad (18)$$

In the mixed strategy Nash Equilibrium, the sensor is indifferent between playing T and NT , i.e., (17) equals (18). Equating (17)

and (18) yields

$$\begin{aligned} & U_k^c(R_{k-1}, R(x_k), \dots, R(x_K)) - U_k^c(R_{k-1}, h, \dots, R(x_K)) \\ &= p_k[U_k^c(R_{k-1}, l, \dots, R(x_K)) - U_k^c(R_{k-1}, h, \dots, R(x_K))]. \end{aligned}$$

Note that when the sensor plays NT , he realizes action $x_k = 0$, therefore, no cost of action is incurred in that case. Based on (13), the equality leads to

$$\Gamma(k)(R(x_k) - h) - g(x_k) = p_k \Gamma(k)(l - h).$$

To make the sensor indifferent between the two actions T and NT , the verification probability should be $p_k = \frac{h - R(x_k) + \frac{g(x_k)}{\Gamma(k)}}{h - l}$.

Now, assume the sensor is mixing between T and NT at stage k with probability q_k . Denote the continuation utility of the operator at stage k by Θ_k ,

$$\Theta_k = \mathbb{E} \left[\sum_{j=k}^K \delta^{j-k} \Pi_j \right] \geq 0. \quad (19)$$

If the operator plays NV at stage k , we denote her continuation utility is by Θ_k^{NV} , and if she plays V , by Θ_k^V . Similar to the sensor's analysis, we start with the last stage. Given Assumption 3, the expected utility of the operator at stage K when she plays V is

$$\begin{aligned} \Theta_K^V &= q_K(S(x_K) - (1 - \omega)R_{K-1} - \omega R(x_K) - C) \\ &\quad + (1 - q_K)(S_K^0 - (1 - \omega)R_{K-1} - \omega l) \\ &= q_K S(x_K) - (1 - \omega)R_{K-1} - q_K \omega R(x_K) \\ &\quad + (1 - q_K)S_K^0 - \omega(1 - q_K)l - C. \end{aligned} \quad (20)$$

The expected utility of the operator at stage K when she plays NV is given by

$$\begin{aligned} \Theta_K^{NV} &= q_N(S(x_K) - (1 - \omega)R_{K-1} - \omega R(x_K)) \\ &\quad + (1 - q_K)(S_K^0 - (1 - \omega)R_{K-1} - \omega h) \\ &= q_K S(x_K) - (1 - \omega)R_{K-1} - q_K \omega R(x_K) \\ &\quad + (1 - q_K)S_K^0 - \omega(1 - q_K)h. \end{aligned}$$

To make the operator indifferent between the two verification decisions, q_K is given by equating Θ_K^V and Θ_K^{NV} , which yields

$$-\omega(1 - q_K)l - C = -\omega(1 - q_K)h \Rightarrow q_K = \frac{h - l - \frac{C}{\omega}}{h - l}.$$

For stage $K - 1$, the expected continuation utility of the operator when playing V is given by

$$\begin{aligned} \Theta_{K-1}^V &= q_{K-1} \left(S(x_{K-1}) - (1 - \omega)R_{K-2} - \omega R(x_{K-1}) \right. \\ &\quad \left. + \delta \Theta_K^V(R(x_{K-1}), R(x_K)) - C \right) \\ &\quad + (1 - q_{K-1}) \left(S_{0_{K-1}} - (1 - \omega)R_{K-2} - \omega l + \delta \Theta_K^V(l, R(x_K)) \right), \end{aligned}$$

the expected continuation utility of the operator when playing NV is given by

$$\begin{aligned} \Theta_{K-1}^{NV} &= q_{K-1} \left(S(x_{K-1}) - (1 - \omega)R_{K-2} - \omega R(x_{K-1}) \right. \\ &\quad \left. + \delta \Theta_K^V(R(x_{K-1}), R(x_K)) - C \right) \\ &\quad + (1 - q_{K-1}) \left(S_{0_{K-1}} - (1 - \omega)R_{K-2} - \omega h + \delta \Theta_K^V(h, R(x_K)) \right). \end{aligned}$$

By equating Θ_{K-1}^V and Θ_{K-1}^{NV} , q_{K-1} can be derived as

$$\begin{aligned} & -\omega(1 - q_{K-1})[l - \delta \Theta_K^V(l, R(x_K))] - C \\ &= -\omega(1 - q_{K-1})[h - \delta \Theta_K^V(h, R(x_K))] \end{aligned}$$

Given the expression for Θ_K^V , we have

$$\Theta_K^V(h, R(x_K)) - \Theta_K^V(l, R(x_K)) = -(1 - \omega)(h - l).$$

Thus, q_{K-1} can be obtained as

$$q_{K-1} = \frac{h - l - \frac{C}{\omega(1 + \delta(1 - \omega))}}{h - l}$$

Similar to derivation of p_k , using backward induction, and given (6) on the reputation score update, we can see that to make the operator indifferent between V and NV , the truth-telling probability by the sensor, q_k , will be

$$q_k = \frac{h - l - \frac{1}{\omega \sum_{j=k}^K \delta^{j-k} (1 - \omega)^{j-k}} C}{h - l}.$$

One can see that the derived mixed strategy NE (p_k, q_k) is unique. Note that for the above mixed strategy to exist, we need to verify that the derived p_k and q_k are valid probabilities. First, note that $q_k \leq 1$ always holds; thus, for q_k to be valid, we require that $\Gamma(k)(h - l) \geq C$ in order to ensure that $q_k \geq 0$. If $\Gamma(k)(h - l) < C$, the sensor will always play NT , in which case the operator should play NV , leading to the operator's outside option. For the operator's side, it is easy to see that $p_k \geq 0$ always holds. Thus, for $p_k \leq 1$ to hold it is required that for each k , $g(x_k) \leq \Gamma(k)(R(x_k) - l)$. Thus, we can summarize the feasibility constraint as

$$g(x_k) \leq \Gamma(k)(R(x_k) - l), \text{ and } \Gamma(k)(h - l) \geq C. \quad (21)$$

which yields

$$\Gamma(k) \geq \max \left\{ \frac{g(x_k)}{R(x_k) - l}, \frac{C}{h - l} \right\}.$$

For these mixed strategies to exist, the operator should select h, l, ω such that at each k (21) is satisfied.

(ii) We first find the optimal strategy chosen by the sensor at stage k , $\tilde{x}_k$. Note that the sensor is a rational decision maker, and therefore chooses his effort level to maximize his expected utility U in (14). Thus, the optimal strategy $\tilde{x}_k$ is given by

$$\frac{\partial \mathbb{E}[U]}{\partial x_k} = \frac{\partial R(x_k)}{\partial x_k} \Gamma(k) - \frac{\partial g(x_k)}{\partial x_k} = 0.$$

Note that concavity of $R(\cdot)$ guarantees the optimality of $\tilde{x}_k$. ■

Proof of Theorem 2: We consider the linear reputations of the form $R(x) = \frac{h-l}{\bar{x}}x + l$. First, using Theorem 1, we conclude that the optimal strategy $\tilde{x}_k$ of the sensor in each stage k is as follows:

$$\tilde{x}_k = \begin{cases} \text{any } x \in [0, \bar{x}] & \text{if } \frac{h-l}{\bar{x}} = \frac{b}{\Gamma(k)} \\ 0 & \text{if } \frac{h-l}{\bar{x}} < \frac{b}{\Gamma(k)}, \\ \bar{x} & \text{if } \frac{h-l}{\bar{x}} > \frac{b}{\Gamma(k)} \end{cases},$$

which means

$$\Gamma(k)R(x_k) - bx_k = \begin{cases} \Gamma(k)l & \text{if } \frac{h-l}{\bar{x}} \leq \frac{b}{\Gamma(k)} \\ \Gamma(k)h - b\bar{x} & \text{if } \frac{h-l}{\bar{x}} > \frac{b}{\Gamma(k)} \end{cases}. \quad (22)$$

Note that when $\frac{h-l}{\bar{x}} = \frac{b}{\Gamma(k)}$ at any stage $1 \leq j < k$ we have $\frac{h-l}{\bar{x}} = \frac{b}{\Gamma(k)} > \frac{b}{\Gamma(j)}$; thus, the sensor will exert effort $\bar{x}$ and receive reputation score h , he will be indifferent about the level of effort in stage k , and exert effort 0 in any stage $k < j \leq K$ since we have $\frac{h-l}{\bar{x}} = \frac{b}{\Gamma(k)} < \frac{b}{\Gamma(j)}$.

Note also when $\frac{h-l}{\bar{x}} > \frac{b}{\Gamma(j)}$, we have $\Gamma(j)h - b\bar{x} \geq \Gamma(j)l$. Given (22) the minimum value of the utility of the sensor at stage k given in (13) will update to

$$\begin{cases} \min\{l, h\}\Delta + \sum_{j=k}^K \delta^{j-k} \Gamma(j)l & \frac{h-l}{\bar{x}} \leq \frac{b}{\Gamma(k)} \\ \min\{l, h\}\Delta + \sum_{j=k}^K \delta^{j-k} \min\{\Gamma(j)h - b\bar{x}, l\}, & \frac{h-l}{\bar{x}} > \frac{b}{\Gamma(k)} \end{cases},$$

where $\Delta = \sum_{j=k}^K (1-\omega)^{j-k+1} \delta^{j-k}$.

As a result, given the last equation, the IR constraint which dictates that $U_k^c \geq 0$ is satisfied regardless of the choice of $h, l \geq 0$. The operator should determine the optimal choice of l to maximize $\mathbb{E}[\Pi]$, subject to the IC constraints of the sensor in each stage, as well as the IR constraint. The operator will thus choose $l = 0$ to provide less payment to the sensor and maximize her own utility. ■

Proof of Theorem 3: We should determine the optimal choice of h, ω to maximize $\mathbb{E}[\Pi]$, subject to the IC constraints of the sensor in each stage, as well as the IR constraint $U_k^c \geq 0$. We therefore substitute for $l = 0$, and rewrite the optimization problem as follows:

$$\max_{\{0 \leq h, 0 \leq \omega \leq 1\}} \mathbb{E}[\Pi] \text{ s.t. IC.}$$

Recall the feasibility constraint for the mixed strategy equilibrium in (21), and replace $R(x_k) = \frac{h}{\bar{x}}x_k, g(x_k) = bx_k$ to obtain

$$b \leq \Gamma(k) \frac{h}{\bar{x}}, \text{ and } \Gamma(k)h \geq C,$$

which yields $h \geq \frac{b\bar{x}}{\Gamma(k)}$, and $h \geq \frac{C}{\Gamma(k)}$. Given that $\Gamma(k+1) < \Gamma(k)$, $\frac{b\bar{x}}{\Gamma(k+1)} > \frac{b\bar{x}}{\Gamma(k)}$, we have the following condition on h chosen by the operator

$$h \geq \frac{b\bar{x}}{\Gamma(K)}, \text{ and } h \geq \frac{C}{\Gamma(K)}.$$

Note that $\Gamma(K) = \omega$. Therefore, for the feasibility constraint to hold, we restrict attention to $h \geq \frac{b\bar{x}}{\omega}$. Note that when $h \geq \frac{b\bar{x}}{\omega}$ the operator will incentivize effort level $\bar{x}$ at the stage $j \neq K$, and the effort x_K^* at the last stage. (Since $h \geq \frac{b\bar{x}}{\omega}$, the sensor is indifferent

about the effort at stage K and we assume he exerts the desired effort by the operator.)

In the following we use backward induction to prove that utility of the operator is decreasing with respect to h when $h \geq \frac{b\bar{x}}{\omega}$. We prove this for Θ_k at any k ; the proof holds for the total utility by setting $k = 1$. We first consider the utility of the operator at stage K

$$\begin{aligned} \Theta_K^{NV} = & -(1-\omega)R_{K-1} + \frac{h-\frac{C}{\omega}}{h} [S(x_K) + \omega(h-R(x_K))] \\ & + \frac{C}{h\omega} S_K^0 - \omega h. \end{aligned}$$

First note that since $\frac{h}{\bar{x}} > \frac{b}{\Gamma(K)}$, we have $x_k = \bar{x}$ for any $k, R_{K-1} = h$ and $R(x_k) = h$. Thus, we can write

$$\begin{aligned} \Theta_K^{NV} = & -(1-\omega)h + \frac{h-\frac{C}{\omega}}{h} \bar{S}_K - \omega h + \frac{C}{h\omega} S_K^0 \\ = & -h + \frac{h-\frac{C}{\omega}}{h} \bar{S}_K + \frac{C}{h\omega} S_K^0. \end{aligned}$$

Now we verify that Θ_K^{NV} is decreasing in h . The derivative of Θ_K^{NV} with respect to h is then given by

$$\frac{\partial \Theta_K^{NV}}{\partial h} = -1 + \frac{C\bar{S}_K}{\omega h^2} - \frac{CS_K^0}{h^2 \omega}.$$

If Assumption 5 holds, i.e., $\sqrt{C(\bar{S}_K - S_K^0)} < b\bar{x}$, we conclude that $C(\bar{S}_K - S_K^0) < \omega h^2$ and consequently $\frac{\partial \Theta_K^{NV}}{\partial h} < 0$. To see this, note that given $b\bar{x} \leq \omega h$ and $0 \leq \omega \leq 1$, we can write

$$C(\bar{S}_K - S_K^0) < (b\bar{x})^2 \leq (\omega h)^2 \leq \omega h^2.$$

Next, we verify that Θ_{K-1}^{NV} is also decreasing in h . Recall that Θ_{K-1}^{NV} can be written as

$$\begin{aligned} \Theta_{K-1}^{NV} = & -(1-\omega)R_{K-2} + q_{K-1}(S(x_{K-1}) \\ & + \Gamma(K-1)(h-R(x_{K-1}))) \\ & + (1-q_{K-1})S_{K-1}^0 + \delta \Theta_K^{NV}(h, x_{K-1}) - \omega h. \end{aligned}$$

Similarly, when $h \geq \frac{b\bar{x}}{\omega}$, the derivative of Θ_{K-1}^{NV} with respect to h is given by

$$\frac{\partial \Theta_{K-1}^{NV}}{\partial h} = -1 + \frac{C\bar{S}_{K-1}}{h^2 \Gamma} - \frac{CS_{K-1}^0}{h^2 \Gamma} + \frac{\partial \Theta_K^{NV}}{\partial h}.$$

Given that Θ_K^{NV} is decreasing in h , it is easy to check that $\frac{\partial \Theta_{K-1}^{NV}}{\partial h}$ is negative and Θ_{K-1}^{NV} is also decreasing in h .

We next consider an arbitrary stage k and we prove Θ_k^{NV} is decreasing in h given that Θ_{k+1}^{NV} is a decreasing function of h . The expected continuation utility of the sensor at stage k is given by

$$\begin{aligned} \Theta_k^{NV} = & -(1-\omega)R_{k-1} + q_k[S(x_{k-1}) \\ & + (h-R(x_{k-1}))\Gamma(k-1)] \\ & - \omega h + \delta \Theta_{k+1}^{NV}(h, \dots, x_k) + \frac{C}{h\Gamma} S_k^0. \end{aligned}$$

When $h \geq \frac{b\bar{x}}{\omega}$, the derivative of Θ_k^{NV} with respect to h is calculated as

$$\frac{\partial \Theta_k^{NV}}{\partial h} = -1 + \frac{C\bar{S}_k}{h^2\Gamma} - \frac{CS_k^0}{h^2\Gamma} + \frac{\partial \Theta_{K+1}^{NV}}{\partial h}.$$

Note that $C(\bar{S}_k - S_k^0) < \Gamma(k)h^2$, since

$$C(\bar{S}_k - S_k^0) < (b\bar{x})^2 \leq (\omega h)^2 \leq \Gamma(K-1)h^2 \leq \dots \leq \Gamma(k)h^2.$$

Thus, assuming Θ_{k+1}^{NV} is a decreasing function of h in conjunction with $C(\bar{S}_k - S_k^0) < \Gamma(k)h^2$ leads to $\frac{\partial \Theta_k^{NV}}{\partial h} < 0$. In other words, we proved that if Θ_{k+1}^{NV} is a decreasing function of h , Θ_k^{NV} is also decreasing in h . Note that the concavity of Θ_k^{NV} at each stage k can be similarly shown using backward induction and the fact that $C(\bar{S}_k - S_k^0) > 0$. Thus, using induction, we conclude that Θ_k^{NV} is decreasing in h for $h \geq \frac{b\bar{x}}{\omega}$ and any k . Therefore, it is optimal for the operator to choose $h = \frac{b\bar{x}}{\omega}$.

The payment at each stage k is thus given by $R(x_k) = \frac{b}{\omega}x_k$. Note that at stage K the operator will therefore incentivize output level x^* which maximizes her own utility at the last stag. Given utility of the operator at the last stage in (20) x^* is given by $\frac{\partial S(x)}{\partial x}|_{x=x^*} = b$. Finally, given the participation constraint of the operator, the sensor will be offered a payment in return for his effort if and only if $\mathbb{E}[\sum_{j=1}^K \delta^{j-1} \Pi_j] \geq 0$. Otherwise, the operator prefers the outside option of not requesting input from the sensor. Part (iii) is proved by substituting the optimal value of the parameters in p_k and q_k in Theorem 1. ■

Proof of Theorem 4: We now analyze the optimal choice of ω given the optimal choice of h identified in Theorem 3. We need to solve the following optimization problem:

$$\max_{\omega} \mathbb{E}[\Pi] = \mathbb{E} \left[\sum_{j=1}^K \delta^{j-1} \Pi_j \right] \geq 0 \quad \text{s.t. } 0 \leq \omega \leq 1.$$

The total utility of the operator over the entire horizon is calculated similar to that of sensor in Lemma 1 and given by

$$\mathbb{E}[\Pi] = \sum_{k=1}^K \delta^{k-1} [q_k(S(x_k) - S_k^0 - \Gamma(k)R(x_k)) + (S_k^0 - C)].$$

If we substitute optimal parameters of the contract, i.e., $l = 0$ and $h = \frac{b\bar{x}}{\omega}$, we can write

$$\begin{aligned} \mathbb{E}[\Pi] &= \sum_{k=1}^K \delta^{k-1} \left[\left(1 - \frac{C\omega}{b\bar{x}\Gamma(k)}\right) (\bar{S}_k - S_k^0 - \Gamma(k)\frac{b\bar{x}}{\omega}) + S_k^0 - C \right] \\ &= \sum_{k=1}^K \delta^{k-1} \left[\bar{S}_k - S_k^0 - \frac{C\omega}{b\bar{x}\Gamma(k)} (\bar{S}_k - S_k^0) + C - \Gamma(k)\frac{b\bar{x}}{\omega} \right] \\ &\quad + \sum_{k=1}^K \delta^{k-1} (S_k^0 - C). \end{aligned}$$

Define $f = \frac{\Gamma(k)}{\omega}$. We take the derivative of the objective function with respect to ω .

$$\frac{\partial \mathbb{E}[\Pi]}{\partial \omega} = - \sum_{k=1}^K \delta^{k-1} b\bar{x} \frac{\partial f(\omega)}{\partial \omega} \left[1 - \frac{1}{f^2} \frac{(\bar{S}_k - S_k^0)C}{(b\bar{x})^2} \right].$$

With the assumption $\sqrt{C(\bar{S}_k - S_k^0)} < b\bar{x}$, and noting that $f(\omega) \geq 1$ and $\frac{\partial f(\omega)}{\partial \omega} < 0$, we conclude that $\frac{\partial \mathbb{E}[\Pi]}{\partial \omega} > 0$. Thus, $\mathbb{E}[\Pi]$ is an increasing function of ω and the optimal choice is to set $\omega = 1$. ■

Proof of Corollary 2: First note that for infinitely repeated games, i.e., when $K \rightarrow \infty$, we have

$$\sum_{j=k}^{\infty} \delta^{j-k} (1-\omega)^{j-k} = \sum_{j'=0}^{\infty} (\delta(1-\omega))^{j'} = \frac{1}{1 - (1-\omega)\delta}.$$

Let Γ denote $\Gamma = \frac{\omega}{1 - (1-\omega)\delta}$. Note that Γ is not a function of k . In this case, p_k and q_k would be updated as

$$p_k = \frac{h - R(x_k) + \frac{g(x_k)}{\Gamma}}{h}, \quad q_k = \frac{h - \frac{1}{\Gamma}C}{h}$$

Thus, the behavior of the sensor becomes more consistent over each stage. Recall that the total utility of the sensor over the entire K stage is given by

$$\mathbb{E}[U] = \sum_{k=1}^K \delta^{k-1} (\Gamma(k)R(x_k) - g(x_k)).$$

We now find the expected payoff of the sensor, and then his optimal action. The total utility of the sensor at stage i is given by

$$U_k^c = (1-\omega)R_{k-1} + \omega R(x_k) - g(x_k) + \delta U_{k+1}^c(\bar{R}_k, x_{k+1}).$$

Notice that

$$\delta U_{k+1}^c(\bar{R}_k, x_{k+1}) = R(x_k) \sum_{j=k+1}^{\infty} \omega \delta^{j-k} (1-\omega)^{j-k} + A,$$

where the terms in A do not include the term $R(x_k)$. Therefore, we can write total utility of the sensor at time i as follows

$$U_k^c = (1-\omega)R_{k-1} - g(x_k) + R(x_k)\Gamma + A$$

Thus, at each stage, the sensor chooses $\bar{x}$ such that

$$\Gamma \frac{\partial R(x_k)}{\partial x_k} = \frac{\partial g(x_k)}{\partial x_k} \bigg|_{x_k=\bar{x}_k}.$$

Next we consider linear payment and effort cost in Assumption 4. Recall the feasibility constraint for the mixed strategy equilibrium and replace $R(x_k) = \frac{h}{\bar{x}}x_k$, $g(x_k) = bx_k$ to obtain

$$b \leq \Gamma \frac{h}{\bar{x}}, \quad \text{and } \Gamma h \geq C.$$

Hence, the operator sets $h \geq \frac{b\bar{x}}{\Gamma}$ which leads to $x_k = \bar{x}$ for any k . Consider the operator's expected continuation utility, when the current reputation of the sensor is R and denote that by Π^R . The expected payoff from verifying and not verifying is the same for the operator at NE, so her expected payoff is equal to her expected payoff from verifying.

At the beginning of the time horizon, we have:

$$\begin{aligned} \Pi^0 &= q(S - \omega R(\bar{x}) + \delta \Pi^{\omega R(\bar{x})} - C) \\ &\quad + (1-q)(S_{01} + \delta \Pi^0 - C) \end{aligned}$$

Given total utility of the operator, we can write

$$\delta \Pi^{\omega R(\bar{x})} - \delta \Pi^0 = -\omega R(\bar{x}) \sum_{k=1}^{\infty} \delta^k (1 - \omega)^k = \Gamma R(\bar{x}).$$

So the utility of the operator can be re-written as

$$\begin{aligned} \Pi^0 &= S - C + \frac{C}{\Gamma h} (S_{0_i} - S) \\ &\quad - \left(1 - \frac{C}{\Gamma h}\right) [\omega R(\bar{x}) - \delta \Pi^{\omega R(\bar{x})}] + \frac{C}{\Gamma h} \delta \Pi^0, \end{aligned}$$

which simplifies to

$$(1 - \delta) \Pi^0 = S - C + \frac{C}{\Gamma h} (S_{0_i} - S) - \left(1 - \frac{C}{\Gamma h}\right) \Gamma R(\bar{x}).$$

So we have

$$\Pi^0 = \frac{1}{1 - \delta} \left[S - C + \frac{C}{\Gamma h} (S_{0_i} - S) - \Gamma h + C \right].$$

We can see that Π^0 is decreasing with respect to h if $b\bar{x} > \sqrt{C(S(\bar{X}) - S_{0_i})}$ since

$$\sqrt{C(S(\bar{X}) - S_{0_i})} < b\bar{x} \leq \Gamma h.$$

Thus, the optimal value of $h = \frac{b\bar{x}}{\Gamma}$. Thus,

$$\Pi^0 = \frac{1}{1 - \delta} \left[S - C + \frac{C}{b\bar{x}} (S_{0_i} - S) - b\bar{x} + C \right].$$

We can see that Π^0 is independent of ω and therefore, the operator is indifferent among all choices of ω . ■

REFERENCES

- [1] H. Gao *et al.*, "A survey of incentive mechanisms for participatory sensing," *IEEE Commun. Surv. Tut.*, vol. 17, no. 2, pp. 918–943, Apr.–Jun. 2015.
- [2] Y. Zhang and M. Van der Schaar, "Collective ratings for online communities with strategic users," *IEEE Trans. Signal Process.*, vol. 62, no. 12, pp. 3069–3083, Jun. 2014.
- [3] W. Saad, Z. Han, H. V. Poor, and T. Basar, "Game-theoretic methods for the smart grid: An overview of microgrid systems, demand-side management, and smart grid communications," *IEEE Signal Process. Mag.*, vol. 29, no. 5, pp. 86–105, Sep. 2012.
- [4] H. He and P. K. Varshney, "A coalitional game for distributed inference in sensor networks with dependent observations," *IEEE Trans. Signal Process.*, vol. 64, no. 7, pp. 1854–1866, Apr. 2016.
- [5] C. G. Cassandras and W. Li, "Sensor networks and cooperative control," *Eur. J. Control*, vol. 11, no. 4, pp. 436–463, 2005.
- [6] R. Borndörfer, J. Buwaya, G. Sagnol, and E. Swarat, "Optimizing toll enforcement in transportation networks: A game-theoretic approach," *Electron. Notes Discrete Math.*, vol. 41, pp. 253–260, 2013.
- [7] M. H. Manshaei, Q. Zhu, T. Alpcan, T. Başar, and J.-P. Hubaux, "Game theory meets network security and privacy," *ACM Comput. Surv.*, vol. 45, no. 3, 2013, Art. no. 25.
- [8] J.-J. Laffont and D. Martimort, *The Theory of Incentives: The Principal-Agent Model*. Princeton, NJ, USA: Princeton Univ. Press, 2009.
- [9] B. Salanié, *The Economics of Contracts: A Primer*. Cambridge, MA, USA: MIT Press, 2005.
- [10] A. Dasgupta and A. Ghosh, "Crowdsourced judgement elicitation with endogenous proficiency," in *Proc. 22nd Int. Conf. World Wide Web*, 2013, pp. 319–330.
- [11] J. A. Bohren and T. Kravitz, "Optimal contracting with costly state verification, with an application to crowdsourcing," 2016. [Online]. Available: https://static1.squarespace.com/static/5a1b26c580bd5e0fc6dbd478/t/5a1c38b5e4966b13422b562e/1511798966391/BohrenKravitz_OptimalContracting_June2016+281
- [12] D. G. Dobakhshari, N. Li, and V. Gupta, "An incentive-based approach to distributed estimation with strategic sensors," in *Proc. IEEE 55th Conf. Decis. Control*, 2016, pp. 6141–6146.
- [13] K. J. Crocker and J. Slemrod, "The economics of earnings manipulation and managerial compensation," *RAND J. Econ.*, vol. 38, no. 3, pp. 698–713, 2007.
- [14] D. Mookherjee and I. Png, "Optimal auditing, insurance, and redistribution," *Quart. J. Econ.*, vol. 104, no. 2, pp. 399–415, 1989.
- [15] K. C. Border and J. Sobel, "Samurai accountant: A theory of auditing and plunder," *Rev. Econ. Studies*, vol. 54, no. 4, pp. 525–540, 1987.
- [16] D. P. Baron and D. Besanko, "Regulation, asymmetric information, and auditing," *RAND J. Econ.*, vol. 15, pp. 447–470, 1984.
- [17] C. Choe, "A mechanism design approach to an optimal contract under ex ante and ex post private information," *Rev. Econ. Des.*, vol. 3, no. 3, pp. 237–255, 1998.
- [18] R. M. Townsend, "Optimal contracts and competitive markets with costly state verification," *J. Econ. Theory*, vol. 21, no. 2, pp. 265–293, 1979.
- [19] N. Miller, P. Resnick, and R. Zeckhauser, "Eliciting informative feedback: The peer-prediction method," *Manage. Sci.*, vol. 51, no. 9, pp. 1359–1373, 2005.
- [20] G. Radanovic and B. Faltings, "Incentive schemes for participatory sensing," in *Proc. Int. Conf. Auton. Agents Multiagent Syst.*, 2015, pp. 1081–1089.
- [21] V. Shnayder, A. Agarwal, R. Frongillo, and D. C. Parkes, "Informed truthfulness in multi-task peer prediction," in *Proc. ACM Conf. Econ. Comput.*, 2016, pp. 179–196.
- [22] G. Radanovic and B. Faltings, "Incentives for subjective evaluations with private beliefs," in *Proc. 29th AAAI Conf. Artif. Intell.*, 2015, pp. 1014–1020.
- [23] X. A. Gao, A. Mao, Y. Chen, and R. P. Adams, "Trick or treat: Putting peer prediction to the test," in *Proc. 15th ACM Conf. Econ. Comput.*, 2014, pp. 507–524.
- [24] F. Farokhi, I. Shames, and M. Cantoni, "Optimal contract design for effort-averse sensors," *Int. J. Cont.*, pp. 1–8, 2018.
- [25] P. Naghizadeh and M. Liu, "Perceptions and truth: A mechanism design approach to crowd-sourcing reputation," *IEEE/ACM Trans. Netw.*, vol. 24, no. 1, pp. 163–176, Feb. 2016.
- [26] G. J. Mailath and L. Samuelson, *Repeated Games and Reputations: Long-Run Relationships*. London, U.K.: Oxford Univ. Press, 2006.
- [27] N. Cao, S. Brahma, and P. K. Varshney, "Target tracking via crowdsourcing: A mechanism design approach," *IEEE Trans. Signal Process.*, vol. 63, no. 6, pp. 1464–1476, Mar. 2015.
- [28] D. Abreu, D. Pearce, and E. Stacchetti, "Toward a theory of discounted repeated games with imperfect monitoring," *Econometrica, J. Econometric Soc.*, vol. 58, pp. 1041–1063, 1990.
- [29] J. A. Bohren, "Stochastic games in continuous time: Persistent actions in long-run relationships," PIER Working Paper, 2014.
- [30] C. Dellarocas, "Reputation mechanism design in online trading environments with pure moral hazard," *Inf. Syst. Res.*, vol. 16, no. 2, pp. 209–230, 2005.
- [31] K. Han, H. Huang, and J. Luo, "Quality-aware pricing for mobile crowdsensing," *IEEE/ACM Trans. Netw.*, vol. 26, no. 4, pp. 1728–1741, Aug. 2018.
- [32] I. Koutsopoulos, "Optimal incentive-driven design of participatory sensing systems," in *Proc. Infocom*, 2013, pp. 1402–1410.
- [33] Y. Zhang and M. Van der Schaar, "Reputation-based incentive protocols in crowdsourcing applications," in *Proc. Infocom*, 2012, pp. 2140–2148.
- [34] J.-S. Lee and B. Hoh, "Sell your experiences: A market mechanism based incentive for participatory sensing," in *Proc. IEEE Int. Conf. Pervasive Comput. Commun.*, 2010, pp. 60–68.
- [35] J.-S. Lee and B. Hoh, "Dynamic pricing incentive for participatory sensing," *Pervasive Mobile Comput.*, vol. 6, no. 6, pp. 693–708, 2010.
- [36] M. O. Jackson, "Mechanism theory," Available at SSRN 2542983, 2003.
- [37] D. G. Dobakhshari, P. Naghizadeh, M. Liu, and V. Gupta, "A reputation-based contract for repeated crowdsensing with costly verification," in *Proc. Amer. Control Conf.*, 2017, pp. 5243–5248.