

Audit Mechanisms for Provable Risk Management and Accountable Data Governance

Jeremiah Blocki, Nicolas Christin, Anupam Datta, and Arunesh Sinha

September 4, 2012

[CMU-CyLab-12-020](#)

[CyLab](#)

Carnegie Mellon University
Pittsburgh, PA 15213

Report Documentation Page		Form Approved OMB No. 0704-0188
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.		
1. REPORT DATE 04 SEP 2012	2. REPORT TYPE	3. DATES COVERED 00-00-2012 to 00-00-2012
4. TITLE AND SUBTITLE Audit Mechanisms for Provable Risk Management and Accountable Data Governance		5a. CONTRACT NUMBER
		5b. GRANT NUMBER
		5c. PROGRAM ELEMENT NUMBER
6. AUTHOR(S)	5d. PROJECT NUMBER	
	5e. TASK NUMBER	
	5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Carnegie Mellon University,CyLab,Pittsburgh,PA,15213		8. PERFORMING ORGANIZATION REPORT NUMBER
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)		10. SPONSOR/MONITOR'S ACRONYM(S)
		11. SPONSOR/MONITOR'S REPORT NUMBER(S)
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited		
13. SUPPLEMENTARY NOTES		
14. ABSTRACT Organizations that collect and use large volumes of personal information are expected under the principle of accountable data governance to take measures to protect data subjects from risks that arise from inappropriate uses of this information. In this paper, we focus on a specific class of mechanisms? audits to identify policy violators coupled with punishments?that organizations such as hospitals, financial institutions, and Web services companies may adopt to protect data subjects from privacy and security risks stemming from inappropriate information use by insiders. We model the interaction between the organization (defender) and an insider (adversary) during the audit process as a repeated game. We then present an audit strategy for the defender. The strategy requires the defender to commit to its action and when paired with the adversary?s best response to it, provably yields an asymmetric subgame perfect equilibrium. We then present two mechanisms for allocating the total audit budget for inspections across all games the organization plays with different insiders. The first mechanism allocates budget to maximize the utility of the organization. Observing that this mechanism protects the organization?s interests but may not protect data subjects, we introduce an accountable data governance property, which requires the organization to conduct thorough audits and impose punishments on violators. The second mechanism we present achieves this property. We provide evidence that a number of parameters in the game model can be estimated from prior empirical studies and suggest specific studies that can help estimate other parameters. Finally, we use our model to predict observed practices in industry (e.g., differences in punishment rates of doctors and nurses for the same violation) and the effectiveness of policy interventions (e.g., data breach notification laws and government audits) in encouraging organizations to adopt accountable data governance practices.		
15. SUBJECT TERMS		

16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 30	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

Audit Mechanisms for Provable Risk Management and Accountable Data Governance^{*}

Jeremiah Blocki, Nicolas Christin, Anupam Datta, and Arunesh Sinha

Carnegie Mellon University, Pittsburgh, PA

{jblocki, nicolasc, danupam, aruneshs}@cmu.edu

Abstract. Organizations that collect and use large volumes of personal information are expected under the principle of *accountable data governance* to take measures to protect data subjects from risks that arise from inappropriate uses of this information. In this paper, we focus on a specific class of mechanisms—audits to identify policy violators coupled with punishments—that organizations such as hospitals, financial institutions, and Web services companies may adopt to protect data subjects from privacy and security risks stemming from inappropriate information use by insiders. We model the interaction between the organization (defender) and an insider (adversary) during the audit process as a repeated game. We then present an audit strategy for the defender. The strategy requires the defender to commit to its action and when paired with the adversary’s best response to it, provably yields an *asymmetric subgame perfect equilibrium*. We then present two mechanisms for allocating the total audit budget for inspections across all games the organization plays with different insiders. The first mechanism allocates budget to maximize the utility of the organization. Observing that this mechanism protects the organization’s interests but may not protect data subjects, we introduce an accountable data governance property, which requires the organization to conduct thorough audits and impose punishments on violators. The second mechanism we present achieves this property. We provide evidence that a number of parameters in the game model can be estimated from prior empirical studies and suggest specific studies that can help estimate other parameters. Finally, we use our model to predict observed practices in industry (e.g., differences in punishment rates of doctors and nurses for the same violation) and the effectiveness of policy interventions (e.g., data breach notification laws and government audits) in encouraging organizations to adopt accountable data governance practices.

^{*} This work was partially supported by the U.S. Army Research Office contract “Perpetually Available and Secure Information Systems” (DAAD19-02-1-0389) to Carnegie Mellon CyLab, the NSF Science and Technology Center TRUST, the NSF CyberTrust grant “Privacy, Compliance and Information Risk in Complex Organizational Processes,” the AFOSR MURI “Collaborative Policies and Assured Information Sharing,” and HHS Grant no. HHS 90TR0003/01. Jeremiah Blocki was also partially supported by a NSF Graduate Fellowship. Arunesh Sinha was also partially supported by the CMU CIT Bertucci Fellowship. The views and conclusions contained in this document are those of the authors and should not be interpreted as representing the official policies, either expressed or implied, of any sponsoring institution, the U.S. government or any other entity.

1 Introduction

Organizations that collect and use large volumes of personal information are expected under the principle of *accountable data governance* to take measures to protect data subjects from risks that arise from these uses of information [1, 2]. In this paper, we focus on a specific class of mechanisms—audits to identify policy violators coupled with punishments—that organizations such as hospitals, financial institutions, and Web services companies may adopt to protect data subjects from privacy and security risks stemming from inappropriate information use by authorized insiders. Indeed, commercial audit tools are emerging to assist in the process of detecting inappropriate information use by insiders [3], and reports of privacy policy violations and associated sanctions are routinely reported in the healthcare sector [4–7].

A central challenge in this setting is the design of effective audit and punishment schemes. We assume that in each audit round audit logs are first analyzed using an automated tool that ranks actions by insiders as potential violations. Our focus is on the next step when a subset of these actions is inspected (because of budgetary constraints) to identify and punish policy violators. We seek to compute the inspection level and punishment level for an “effective” scheme.

The challenge in modeling the complex interaction between the auditor and audited agent includes making reasonable abstractions and assumptions. We model the interaction between an organization (the defender) and the insider (the adversary) as a repeated game with imperfect information (the defender does not observe the adversary’s actions) and public signals (the outcome of the audit is public). The model captures a number of important economic considerations that influence the design of audit mechanisms. The game model (described in Section 3) replaces the byzantine adversary model in our previous work [8] with a *near-rational adversary model*. These adversaries act rationally with high probability and in a byzantine manner otherwise (similar to a *trembling hand* assumption [9]). Adversaries benefit from violations they commit (e.g., by selling personal data) and suffer due to punishments imposed for detected violations. The model generalizes from the situation in which the defender interacts with a single adversary to one where she interacts with multiple, non-colluding adversaries via a natural product game construction that we define. Each audit game is parametrized by a budget that the defender can use to conduct inspections.

We then present an audit strategy for the defender. This strategy when paired with the adversary’s best response to it provably yields an *asymmetric approximate subgame perfect equilibrium* (Theorem 1). This equilibrium concept implies that the adversary does not gain at all from deviating from her best response strategy (see Section 4). We define this equilibrium concept by adapting the standard notion of approximate subgame perfect equilibrium, which has a symmetric flavor and permits both players to obtain small gains by unilaterally deviating from their equilibrium strategy. The symmetric equilibrium concept is unsuitable for our security application, where an adversary who deviates motivated by a small gain could cause a big loss for the organization. The defender’s strategy involves committing to a level of inspection and punishment. The strategy has two desirable properties. First, the commitment results in a predictable equilibrium since the adversary plays her best response to the strategy. Second, the strategy is *deterrence dominant* over the set of maximum utility defender strategies that

result in a perfect public equilibrium, i.e., whenever such a strategy deters the adversary, so does our audit strategy (see Theorem 2 for the formal statement).

We design two mechanisms using which the defender can allocate her total audit budget across the different games to audit different insiders and types of potential violations. The first mechanism optimizes the defender’s utility. Observing that this mechanism protects the organization’s interests but may not protect data subjects, we introduce an accountable data governance property, which places an operational requirement on the organization to use a sufficiently effective log analysis tool and maintain sufficiently high inspection and punishment rates. The second mechanism allocates the total audit budget to achieve this property (see Section 5).

Finally, we demonstrate the usefulness of our model by predicting and explaining observed practices in industry (e.g., differences in punishment rates of doctors and nurses for the same violation) and analyzing the effectiveness of policy interventions (e.g., data breach notification laws and government audits) in encouraging organizations to adopt accountable data governance practices (see Section 6). We present comparisons to additional related work in Section 7 and conclusions and directions for future work in Section 8.

2 Overview

In this section, we provide an overview of our model using a motivating scenario that will serve as a running example for this paper. Consider a “Hospital X” with employees in different roles (doctors, nurses). X conducts weekly audits to ensure that accesses to personal health records are legitimate. Given budget constraints, X cannot check every single access. The first step in the audit process is to analyze the access logs using an automated tool that ranks accesses as potential violations. Hospital X assesses the (monetary) impact of different types of violations and decides what subset to focus on by balancing the cost of audit and the expected impact (“risk”) from policy violations. This type of audit mechanism is common in practice [10–13].

We provide a game model for this audit process. An employee (“adversary,” $\mathcal{A}$) executes tasks, i.e., actions that are permitted as part of their job. We only consider tasks that can later be audited, e.g., through inspection of logs. For example, in X the tasks are accesses to health records. We can distinguish $\mathcal{A}$ ’s tasks between legitimate tasks and violations of a policy. Different *types of violations* may have different impact on the organization. We assume that there are K different types of violations that $\mathcal{A}$ can commit. Examples of violations of different types in Hospital X include inappropriate access to a celebrity’s health record, or access to a health record leading to identity theft. $\mathcal{A}$ benefits by committing violations: the benefit is quantifiable using information from existing studies or by human judgment. For example, reports [14, 15] indicate that on average the *personal benefit* of a hospital employee from selling a common person’s health record is \$50. On the other hand, if $\mathcal{A}$ is caught committing a violation then she is punished according to the *punishment policy* used by $\mathcal{D}$. For example, employees could be terminated, as happened in similar recent incidents [6, 7].

The organization $\mathcal{D}$ can classify each adversary’s task by type. However, $\mathcal{D}$ cannot determine with certainty whether a particular task is legitimate or a violation without

investigating. Furthermore, $\mathcal{D}$ cannot inspect all of $\mathcal{A}$'s tasks due to *budgetary constraints*. As such, some violations may go undetected *internally*, but could be detected *externally*. Governmental audits, whistle-blowing, patient complaints [16, 17] are all examples of situations that could lead to external detection of violations. Externally detected violations usually cause more economic damage to the organization than internally caught violations. The 2011 Ponemon Institute report [18] states that patients whose privacy has been violated are more likely to leave (and possibly sue) a hospital if they discover the violation on their own than if the hospital detects the violation and proactively notifies the patient.

The economic impact of a violation is a combination of *direct and indirect costs*; direct costs include breach notification and remedial cost, and indirect costs include loss of customers and brand value. For example, the 2010 Ponemon Institute report [19] states that the average cost of privacy breach *per record* in health care is \$301 with indirect costs about two thirds of that amount. Of course, certain violations may result in much higher direct costs, e.g., \$25,000 per record (up to \$250,000 in total) in fines alone in the state of California [6]. These fines may incentivize organizations to adopt aggressive punishments policies. However, severe punishment policies create a hostile work environment resulting in economic losses for the organization due to low employee motivation and a failure to attract new talent [20].

The organization needs to balance auditing costs, potential economic damages due to violations and the economic impact of the punishment policy. The employees need to weigh their gain from violating policies against loss from getting caught by an audit and punished. The actions of one party impact the actions of the other party: if employees never violate, the organization does not need to audit; likewise, if the organization never audits, employees can violate policies in total impunity. Given this strategic interdependency, we model the auditing process as a *repeated game* between the organization and its employees, where the discrete rounds characterize audit cycles. The game is parameterized by quantifiable variables such as the personal benefit of employee, the cost of breach, and the cost of auditing, among others. The organization is engaged in multiple such games simultaneously with different employees and has to effectively allocate its total audit budget across the different games.

3 Audit Game Model

We begin by providing a high level view of the audit process, before describing the audit game in detail (Section 3). In practice, the organization is not playing a repeated audit game against a specific employee, but against all of its n employees at the same time. However, if we assume that 1) a given employee's actions for a type of task are independent of her actions for other types, and that 2) employees do not collude with other employees and act independently, we can decompose the overall game into nK independent *base* repeated games, that the organization plays in parallel. One base repeated game corresponds to a given type of access k by a given employee $\mathcal{A}$, and will be denoted by $\mathcal{G}_{\mathcal{A},k}$. Each game $\mathcal{G}_{\mathcal{A},k}$ is described using many parameters, e.g., *loss due to violations*, *personal benefit* for employee, etc. We abuse notation in using $\mathcal{G}_{\mathcal{A},k}$ to refer to a base repeated game of type k with any value of the parameters.

In our proposed audit process the organization follows the steps below in each audit cycle for every game $\mathcal{G}_{\mathcal{A},k}$. Assume the parameters of the game have been estimated and the equilibrium audit strategy computed for the first time auditing is performed.

before audit:

1. If any parameter changes go to step 2 else go to *audit*.
2. Estimate parameters. Compute equilibrium of $\mathcal{G}_{\mathcal{A},k}$.

audit:

3. Audit using actions of the computed equilibrium.

Note that the parameters of $\mathcal{G}_{\mathcal{A},k}$ may change for any given round of the game, resulting in a different game. However, neither $\mathcal{D}$ nor $\mathcal{A}$ knows when that will happen. As such, since the horizon of $\mathcal{G}_{\mathcal{A},k}$ with a fixed set of parameters is infinite, we can describe the interaction between the organization and its employees with an infinitely repeated game for the period in which the parameters are unchanged (see [9] for details). Thus, the game $\mathcal{G}_{\mathcal{A},k}$ is an infinitely repeated game of *imperfect information* since $\mathcal{A}$'s action is not directly observed. Instead, noisy information about the action, called a *public signal* is observed. The public signal here consists of a) the detected violations b) number of tasks by $\mathcal{A}$ and c) $\mathcal{D}$'s action. The K parallel games played between $\mathcal{A}$ and $\mathcal{D}$ can be composed in a natural manner into one repeated game (which we call $\mathcal{G}_{\mathcal{A}}$) by taking the product of action spaces and adding up utilities from the games.

Finally, analyzing data to detect changes of parameters may require the use of statistical methods [21], data mining and learning techniques. We do not delve into details of these methods as that is beyond the scope of this paper and estimating risk parameters has been studied extensively in many contexts [10–13, 15]. Observe that change of parameters may change the equilibrium of the game, e.g., a lot of violations in quick succession by an employee (in spite of being inspected sufficiently) may result in the organization changing the personal benefit of the employee leading to more inspection.

Formal Description In the remainder of this section, we focus on the base repeated games $\mathcal{G}_{\mathcal{A},k}$. We use the following notations in this paper:

- Vectors are represented with an arrow on top, e.g., $\vec{v}$ is a vector. The i^{th} component of a vector is given by $\vec{v}(i)$. $\vec{v} \leq \vec{a}$ means that both vectors have the same number of components and for any component i , $\vec{v}(i) \leq \vec{a}(i)$.
- Random variables are represented in boldface, e.g., $\mathbf{x}$ and $\mathbf{X}$ are random variables.
- $E(\mathbf{X})[q, r]$ denotes the expected value of random variable X , when particular parameters of the probability mass function of $\mathbf{X}$ are set to q and r .
- We will use a shorthand form by dropping $\mathcal{A}$, k and the vector notation, as we assume these are implicitly understood for the game $\mathcal{G}_{\mathcal{A},k}$, i.e., a quantity $\vec{x}_{\mathcal{A}}(k)$ will be simply denoted as x . We use this form whenever the context is restricted to game $\mathcal{G}_{\mathcal{A},k}$ only.

$\mathcal{G}_{\mathcal{A},k}$ is fully defined by the players, the time granularity at which the game is played, the actions the players can take, and the utility the players obtain as a result of the actions they take. We next discuss these different concepts in turn.

Players: The game $\mathcal{G}_{\mathcal{A},k}$ is played between the organization $\mathcal{D}$ and an adversary $\mathcal{A}$. For instance, the players are hospital $\mathbf{X}$ and a nurse in $\mathbf{X}$.

Round of play: In practice, audits for all employees and all types of access are performed together and usually periodically. Thus, we adopt a discrete-time model, where

time points are associated with rounds. Each round of play corresponds to an audit cycle. We group together all of the $\mathcal{A}$'s actions (tasks of a given type) in a given round. All games $\mathcal{G}_{\mathcal{A},k}$ are synchronized, i.e., all rounds t in all games are played simultaneously.

Adversary action space: In each round, the adversary $\mathcal{A}$ chooses two quantities of type k : the number of tasks she performs, and the number of such tasks that are violations. If we denote by U_k the maximum number of type k tasks that any employee can perform, then $\mathcal{A}$'s entire action space for $\mathcal{G}_{\mathcal{A},k}$ is given by $A_k \times V_k$ with $A_k = \{u_k, \dots, U_k\}$ ($u_k \leq U_k$) and $V_k = \{1, \dots, U_k\}$. Let $\vec{a}_{\mathcal{A}}^t$ and $\vec{v}_{\mathcal{A}}^t$ be vectors of length K such that the components of vector $\vec{a}$ are the number of tasks of each type that $\mathcal{A}$ performs at time t , and the components of vector $\vec{v}$ are the number of violations of each type. Since violations are a subset of all tasks, we always have $\vec{v}_{\mathcal{A}}^t \leq \vec{a}_{\mathcal{A}}^t$. In a given audit cycle, $\mathcal{A}$'s action in the game $\mathcal{G}_{\mathcal{A},k}$ is defined by $\langle \vec{a}_{\mathcal{A}}^t(k), \vec{v}_{\mathcal{A}}^t(k) \rangle$, that is $\langle a^t, v^t \rangle$ in shorthand form, with $a^t \in A_k$ and $v^t \in V_k$.

Instead of being perfectly rational, we model $\mathcal{A}$ as playing with a *trembling hand* [9]. Whenever $\mathcal{A}$ chooses to commit v^t violations in as given round t , she does so with probability $1 - \epsilon_{th}$, but, with (small) probability ϵ_{th} she commits some other number of violations sampled from an unknown distribution D_0^t over all possible violations. In other words, we allow $\mathcal{A}$ to act completely arbitrarily when she makes a mistake. For instance, a nurse in X may lose her laptop containing health records leading to a breach.

Defender action space: $\mathcal{D}$ also chooses two quantities of type k in each round: the number of inspections to perform, and the punishment to levy for each type- k violation detected. Let $\vec{s}_{\mathcal{A}}^t$ be the vector of length K such that components of vector $\vec{s}_{\mathcal{A}}^t$ are the number of inspections of each type that $\mathcal{D}$ performs in round t . The number of inspections that $\mathcal{D}$ can conduct is bounded by the number of tasks that $\mathcal{A}$ performs, and thus, $\vec{s}_{\mathcal{A}}^t \leq \vec{a}_{\mathcal{A}}^t$. $\mathcal{D}$ uses a log analysis tool $\mathcal{M}$ to sort accesses according to the probability of them being a violation. Then, $\mathcal{D}$ chooses the top $\vec{s}_{\mathcal{A}}^t(k) = s^t$ tasks from the sorted output of $\mathcal{M}$ to inspect in game $\mathcal{G}_{\mathcal{A},k}$. Inspection is assumed perfect, i.e., if a violation is inspected, it is detected. The number of inspections is bounded by budgetary constraints. Denoting the functions that outputs cost of inspection for each type of violation by $\vec{C}$, we have $\vec{C}(k)(\vec{s}_{\mathcal{A}}^t(k)) \leq \vec{b}_{\mathcal{A}}^t(k)$ where $\vec{b}_{\mathcal{A}}^t(k)$ defines a per-employee, per-type budget constraint. The budget allocation problem is an optimization problem depending on the audit strategy, which we discuss in Section 5.1.

$\mathcal{D}$ also chooses a punishment rate $\vec{P}_{\mathcal{A}}^t(k) = P^t$ (fine per violation of type k) in each round t to punish $\mathcal{A}$ if violations of type k are detected. P^t is bounded by a maximum punishment P_f corresponding to the employee being fired, and the game terminated.

Finally, $\mathcal{D}$'s choice of the inspection action can depend only on $\mathcal{A}$'s total number of tasks, since the number of violations is not observed. Thus, $\mathcal{D}$ can choose its strategy as a function from number of tasks to inspections and punishment even before $\mathcal{A}$ performs its action. In fact, we simulate $\mathcal{D}$ acting first and the actions are *observable* by requiring $\mathcal{D}$ to commit to a strategy and provide a proof of honoring the commitment. Specifically, $\mathcal{D}$ computes its strategy, makes it public and provides a proof of following the strategy after auditing is done. The proof can be provided by maintaining an audit trail of the audit process itself.

Outcomes: We define the outcome of a single round of $\mathcal{G}_{\mathcal{A},k}$ as the number of violations detected in internal audit and the number of violations detected externally. We

assume that there is a fixed exogenous probability p ($0 < p < 1$) of an internally undetected violation getting caught externally. Due to the probabilistic nature of all quantities, the outcome is a random variable. Let $\vec{\mathbf{O}}_{\mathcal{A}}^t$ be the vector of length K such that the $\vec{\mathbf{O}}_{\mathcal{A}}^t(k) = \mathbf{O}^t$ represents the outcome for the t^{th} round for the game $\mathcal{G}_{\mathcal{A},k}$. Then $\mathbf{O}^t$ is a tuple $\langle \mathbf{O}_{int}^t, \mathbf{O}_{ext}^t \rangle$ of violations caught internally and externally. As stated earlier, we assume the use of a log analysis tool $\mathcal{M}$ to rank the accesses with more likely violations being ranked higher. Then, the probability mass function for $\vec{\mathbf{O}}_{int}^t$ is a distribution parameterized by $\langle a^t, v^t \rangle$, s and $\mathcal{M}$. The baseline performance of $\mathcal{M}$ is when the s accesses to be inspected are chosen at random, resulting in a hyper-geometric distribution with mean $v^t \alpha^t$, where $\alpha^t = s^t / a^t$. We assume that the mean of the distribution is $\mu(\alpha^t) v^t \alpha^t$, where $\mu(\alpha^t)$ is a function dependent on α^t that measures the *performance* of $\mathcal{M}$ and $\forall \alpha^t \in [0, 1]. \mu \geq \mu(\alpha^t) \geq 1$ for some constant μ (μ is overloaded here). Note that we must have $\mu(\alpha^t) \alpha^t \leq 1$, and further, we assume that $\mu(\alpha^t)$ is monotonically non-increasing in α^t . The probability mass function for $\mathbf{O}_{ext}^t$ conditioned on $\mathbf{O}_{int}^t$ is a binomial distribution parameterized by p .

Utility functions: In a public signaling game like $\mathcal{G}_{\mathcal{A},k}$, the utilities of the players depend only on the public signal and their own action, while the strategies they choose depend on the history of public signals [22]. The utility of the repeated game is defined as a (delta-discounted) sum of the expected utilities received in each round, where the expectation is taken with respect to the distribution over histories. Let the discount factor for $\mathcal{D}$ be $\delta_{\mathcal{D}}$ and for any employee $\mathcal{A}$ be $\delta_{\mathcal{A}}$. We assume that $\mathcal{D}$ is patient, i.e., future rewards are almost as important as immediate rewards, and $\delta_{\mathcal{D}}$ is close to 1. $\mathcal{A}$ is less patient than $\mathcal{D}$ and hence $\delta_{\mathcal{A}} < \delta_{\mathcal{D}}$.

Defender utility function: $\mathcal{D}$'s utility in a round of the game $\mathcal{G}_{\mathcal{A},k}$ consists of the sum of the *cost of inspecting* $\mathcal{A}$'s actions, the *monetary loss from a high punishment rate* for $\mathcal{A}$, and *direct and indirect costs* of violations. As discussed before, inspection costs are given by $C(s^t)$ where $C = \vec{C}(k)$ is a function denoting the cost of inspecting type- k tasks. Similarly, the monetary loss from losing employee's productivity due to fear of punishment is given by $e(P^t)$, where $e = \vec{e}_A(k)$ is a function for type- k tasks. The functions in $\vec{C}$ and $\vec{e}$ must satisfy the following constraints: 1) they should be monotonically increasing in the argument and 2) $\vec{C}(k) \geq 0$, $\vec{e}_A(k) \geq 0$ for all k .

We characterize the effect of violations on the organization's indirect cost similarly to the reputation loss as in previous work [8]. Additionally, the generic function described below is capable of capturing direct costs, as shown in the example following the function specification. Specifically, we define a function r_k (r in shorthand form) that, at time t , takes as input the number of type- k violations caught internally, the number of type- k violations caught externally, and a time horizon τ , and outputs the overall loss at time $t + \tau$ due to these violations at time t . r is stationary (i.e., independent of t), and externally caught violations have a stronger impact on r than internally detected violations. Further, $r(\langle 0, 0 \rangle, \tau) = 0$ for any τ (undetected violations have 0 cost), and r is monotonically decreasing in τ and becomes equal to zero for $\tau \geq m$ (violations are forgotten after a finite amount of rounds). As in previous work [8], we construct the utility function at round t by immediately accounting for future losses due to violations occurring at time t . This allows us to use standard game-theory results, while at the same time, providing a close approximation of the defender's loss [8]. With these

notations, $\mathcal{D}$'s utility at time t in $\mathcal{G}_{\mathcal{A},k}$ is

$$\mathbf{Rew}_{\mathcal{D}}^t(\langle s^t, P^t \rangle, \mathbf{O}^t) = - \sum_{j=0}^{m-1} \delta_{\mathcal{D}}^j r(\mathbf{O}^t, j) - C(s^t) - e(P^t). \quad (1)$$

This per-round utility is always negative (or at most zero). As is typical of security games (e.g., [23, 24] and related work), implementing security measures does not provide direct benefits to the defender, but is necessary to pare possible losses. Hence, the goal for the defender is to have this utility as close to zero as possible.

The above function can capture direct costs of violations as an additive term at time $\tau = 0$. As a simple example [8], assuming the average direct costs for internally and externally caught violations are given by R_{int}^D and R_{ext}^D , and the function r is linear in the random variables $\vec{\mathbf{O}}_{int}^t$ and $\vec{\mathbf{O}}_{ext}^t$, r can be given by

$$r(\mathbf{O}^t, \tau) = \begin{cases} (c + R_{int}^D) \mathbf{O}_{int}^t + (\psi c + R_{ext}^D) \mathbf{O}_{ext}^t & \text{for } \tau = 0 \\ \delta^\tau c (\mathbf{O}_{int}^t + \psi \cdot \mathbf{O}_{ext}^t) & \text{for } 1 \leq \tau < m \\ 0 & \text{for } \tau \geq m, \end{cases}$$

where $\delta \in (0, 1)$ and $\psi \geq 1$. Then Eqn. (1) reduces to

$$\mathbf{Rew}_{\mathcal{D}}^t(\langle s^t, P^t \rangle, \mathbf{O}^t) = -R_{int} \mathbf{O}_{int}^t - R_{ext} \mathbf{O}_{ext}^t - C(s^t) - e(P^t), \quad (2)$$

with $R_{int} = R_{int}^I + R_{int}^D$, $R_{int}^I = c(1 - \delta^m \delta_{\mathcal{D}}^m)/(1 - \delta \delta_{\mathcal{D}})$ and $R_{ext} = \psi R_{int}^I + R_{ext}^D$.

Adversary utility function: We define $\mathcal{A}$'s utility as the sum of $\mathcal{A}$'s *personal benefit* gained by committing violations and the *punishment* that results due to detected violations. Personal benefit is a monetary measure of the benefit that $\mathcal{A}$ gets out of violations. It includes all kinds of benefits, e.g., curiosity, actual monetary benefit (by selling private data), revenge, etc. It is natural that true personal benefit of $\mathcal{A}$ is only known to $\mathcal{A}$. Our model of personal benefit of $\mathcal{A}$ is linear and is defined by a rate of personal benefit for each type of violation given by the vector $\vec{I}_{\mathcal{A}}$ of length K . The punishment is the vector $\vec{P}_{\mathcal{A}}^t$ of length K chosen by $\mathcal{D}$, as discussed above. Using shorthand notation, $\mathcal{A}$'s utility, for the game $\mathcal{G}_{\mathcal{A},k}$, is:

$$\mathbf{Rew}_{\mathcal{A}}^t(\langle a^t, v^t \rangle, \langle s^t, P^t \rangle, \mathbf{O}^t) = Iv^t - P^t (\mathbf{O}_{int}^t + \mathbf{O}_{ext}^t).$$

Observe that the utility function of a player depends on the public signal (observed violations, $\mathcal{D}$'s action) and the action of the player, which conforms to the definition of a repeated game with imperfect information and public signaling. In such games, the *expected* utility is used in computing equilibria.

Let $\alpha^t = s^t/a^t$ and $\nu(\alpha^t) = \mu(\alpha^t)\alpha^t$. Then, $E(\mathbf{O}_{int}^t) = \nu(\alpha^t)v^t$, and $E(\mathbf{O}_{ext}^t) = p\nu^t(1 - \nu(\alpha^t))$. The expected utilities in each round then become:

$$\begin{aligned} E(\mathbf{Rew}_{\mathcal{D}}^t) &= - \sum_{j=0}^{m-1} \delta_{\mathcal{D}}^j E(r(\mathbf{O}^t, j)) [v^t, a^t, \alpha^t] - C(\alpha^t a^t) - e(P^t), \\ E(\mathbf{Rew}_{\mathcal{A}}^t) &= Iv^t - P^t v^t (\nu(\alpha^t) + p(1 - \nu(\alpha^t))). \end{aligned}$$

The expected utility of $\mathcal{A}$ depends only on the level of inspection and not on the actual number of inspections. For the example loss function given by Eqn. (2), the utility function of $\mathcal{D}$ becomes:

$$E(\mathbf{Rew}_{\mathcal{D}}^t) = -v^t(R_{int}\nu(\alpha^t) + R_{ext}p(1 - \nu(\alpha^t))) - C(\alpha^t a^t) - e(P^t).$$

In addition to the action dependent utilities above, the players also receive a fixed utility every round, which is the salary for $\mathcal{A}$ and value generated by $\mathcal{A}$ for $\mathcal{D}$. P_f depends on these values, and is calculated in Appendix B.2. Finally, the model parameters that may change over time are R_{ext} , R_{int} , p , function C , function e , function μ and I .

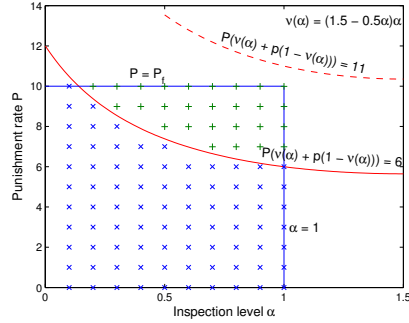


Fig. 1. Non-deterred ($\times$) and deterred ($+$) region for $I = \$6$. $I = \$11$ has empty deterred region.

Graphical representation: A graphical representation of the utilities helps illustrate the ideas presented in the next two sections. (See Figure 1). Consider the 2-dimensional plane $R^{\alpha, P}$ spanned by α^t and P^t . We define a feasible audit space in $R^{\alpha, P}$ given by $0 \leq \alpha^t \leq 1$ and $0 \leq P^t \leq P_f$. $\mathcal{D}$'s actions are points in the feasible region. The expected utility of the adversary in each round is given by $v^t(I - P^t(\nu(\alpha^t) + p(1 - \nu(\alpha^t))))$. Thus, the curve in $R^{\alpha, P}$ given by $I = P^t(\nu(\alpha^t) + p(1 - \nu(\alpha^t)))$ is the separator between positive and negative expected utility regions for the adversary in each round. Within the feasible region, we call the region of positive expected utility the *non-deterred region* and the region of negative utility the *deterred region*.

$\mathcal{A}$'s utility can as well be non-linear, e.g., if $\mathcal{D}$ decides to scale punishment quadratically with violations. Technically, this partitions the feasible audit space into many regions, with each region associated with the number of violations that maximize the utility of $\mathcal{A}$ in that region. We emphasize that the equilibrium presented later can be easily extended to consider such cases. To keep the presentation simple we keep using the linear utility throughout the paper, which yields two regions associated with 0 or all violations. Similarly, it is possible to add any other relevant term to $\mathcal{D}$'s utility, e.g., if $\mathcal{D}$ satisfies a certain accountability criteria (defined later in Section 5) then it may earn positive benefit out of increased reputation.

Estimation: Next, we describe techniques of estimating parameters of game $\mathcal{G}_{\mathcal{A}, k}$, obtaining sample estimates in the process. Before getting to constant values, we state the functions that we use as concrete instances for the examples in this paper. We use simple linear functions for audit cost ($C(\alpha a) = C\alpha a$) and for punishment loss ($e(P) = eP$). The performance of $\mathcal{M}$ is dependent on the tool being used and we use a linear function for $\mu(\cdot)$ to get $\nu(\alpha) = \mu\alpha - (\mu - 1)\alpha^2$, where μ is a constant. Further, we use the example loss function (with R_{int} and R_{ext}) stated in the last sub-section. We note that our theorems work with any function; these functions above are the simplest functions that satisfy the constraints on these functions stated in the last sub-section. Next, we gather data from industry wide studies to obtain sample estimates for parameters.

As stated in Section 2, values of direct and indirect costs of violation (average of R_{int} and R_{ext} is \$300 in healthcare [19], a detailed breakdown is present in the ANSI report [15]), maximum personal benefit I (\$50 for medical records [14, 15]), etc. are available in studies. Also, in absence of studies quantitatively distinguishing externally and internally caught violations we assume $R_{int} = R_{ext} = \$300$. Many parameters depends on the employee, his role in the organization and type of violation. Keeping a track of violations and behavior within the organization offers a data source for estimating and detecting changes in these parameters. We choose values for these parameters that are not extremes, $e = \$10$, $I = \$6$, $\epsilon_{th} = 0.03$, $\delta_A = 0.4$ and $U_k = 40$. Further, under certain assumptions we calculate P_f (in Appendix B.2) to get $P_f = \$10$. Finally, the average cost of auditing C and performance factor μ of log analysis tool should be known to $\mathcal{D}$. We assume values $C = \$50$, and tool performance $\mu = 1.5$.

4 Auditing Strategy

In this section, we define a suitable equilibrium concept for the audit game (Section 4.1) and present a strategy for the defender such that the best response to that strategy by the adversary results in an equilibrium being attained (Section 4.2). Finally, we compare our equilibrium with other equilibria (Section 4.3). Recall that the equilibrium of the game occurs in the period in which the game parameters are fixed.

4.1 Equilibrium Concepts

We begin by introducing standard terminology from game theory. In a one-shot extensive form game players move in order. We assume player 1 moves first followed by player 2. An extensive form repeated game is one in which the round game is a one-shot extensive game. The history is a sequence of actions. Let H be the set of all possible histories. Let S_i be the action space of player i . A strategy of player i is a function $\sigma_i : H_i \rightarrow S_i$, where $H_i \subset H$ are the histories in which player i moves. The utility in each round is given by $r_i : S_1 \times S_2 \rightarrow \mathbb{R}$. The total utility is a δ_i -discounted sum of utilities of each round, normalized by $1 - \delta_i$.

The definition of strategies extends to extensive form repeated games with public signals. We consider a special case here that resembles our audit game. Player 1 moves first and the action is observed by player 2, then player 2 moves, but, that action may not be perfectly observed, instead resulting in a public signal. Let the space of public signals be Y . In any round, the observed public signal is distributed according to the distribution $\Delta Y(\cdot|s)$, i.e., $\Delta Y(y|s)$ is the probability of seeing signal y when the action profile s is played. In these games, a history is defined as an alternating sequence of player 1's action and public signals, ending in a public signal for histories in which player 1 has to move and ending in player 1's move for histories in which player 2 has to move. The actual utility in each round is given by the function $r_i : S_i \times Y \rightarrow \mathbb{R}$. The total expected utility g_i is the expected normalized δ_i -discounted sum of utilities of each round, where the expectation is taken over the distribution over public signals and histories. For any history h , the game to be played in the future after h is called the *continuation game* of h with total utility given by $g_i(\sigma, h)$.

A strategy profile (σ_1, σ_2) is a *subgame perfect equilibrium* (SPE) of a repeated game if it is a Nash equilibrium for all continuation games given by any history h [9]. One way of determining if a strategy is a SPE is to determine whether the strategy satisfies the *single stage deviation* property, that is, any *unilateral deviation* by any player in any single round is not profitable. We define a natural extension of SPE, which we call *asymmetric subgame perfect equilibrium* (or (ϵ_1, ϵ_2) -SPE), which encompasses SPE as a special case when $\epsilon_1 = \epsilon_2 = 0$.

Definition 1. $((\epsilon_1, \epsilon_2)$ -SPE) Denote concatenation operator for histories as $;$. Strategy profile σ is a (ϵ_1, ϵ_2) -SPE if for history h in which player 1 has to play, given $h' = h; \sigma_1(h)$ and $h'' = h; s_1$,

$$\begin{aligned} & E(r_1(\sigma_1(h), \mathbf{y}))[\sigma_1(h), \sigma_2(h')] + \delta_1 E(g_1(\sigma, h'; \mathbf{y}))[\sigma_1(h), \sigma_2(h')] \\ & \geq E(r_1(s_1, \mathbf{y}))[\sigma_1(h), \sigma_2(h'')] + \delta_1 E(g_1(\sigma, h''; \mathbf{y}))[\sigma_1(h), \sigma_2(h'')] - \epsilon_1 \end{aligned}$$

for all s_1 . For history h in which player 2 has to play, given $a(h)$ is the last action by player 1 in h , for all s_2

$$\begin{aligned} & E(r_2(\sigma_2(h), \mathbf{y}))[\sigma_1(h), \sigma_2(h)] + \delta_2 E(g_2(\sigma, h; \mathbf{y}))[\sigma_1(h), \sigma_2(h)] \\ & \geq E(r_2(s_2, \mathbf{y}))[\sigma_1(h), \sigma_2(h)] + \delta_2 E(g_2(\sigma, h; \mathbf{y}))[\sigma_1(h), \sigma_2(h)] - \epsilon_2 \end{aligned}$$

We are particularly interested in $(\epsilon_1, 0)$ -SPE, where player 1 is the defender and player 2 is the adversary. By setting $\epsilon_2 = 0$, we ensure that a rational adversary will never deviate from the expected equilibrium behavior. Such equilibria are important in security games, since $\epsilon_2 > 0$ could incentivize the adversary to deviate from her strategy, possibly resulting in significant loss to the defender.

The following useful property about history-independent strategies, which follows directly from the definition, helps in understanding our proposed history-independent audit strategy.

Property 1. If a strategy profile σ is history-independent, i.e., $\sigma_1(h) = \sigma_1()$ and $\sigma_2(h) = \sigma_2(a(h))$ then the condition to test for SPE reduces to $E(r_1(\sigma_1(), \mathbf{y})) \geq E(r_1(s_1, \mathbf{y}))$, for player 1 and to $E(r_2(\sigma_2(h), \mathbf{y})) \geq E(r_2(s_2, \mathbf{y}))$, for player 2, since $g_i(\sigma, h; y)$ is the same for all y and each i . Also, if $E(r_i(s_i, \mathbf{y})) - E(r_i(\sigma_i(h), \mathbf{y})) \leq \epsilon_i$ for all i and s_i then σ is an (ϵ_1, ϵ_2) -SPE strategy profile.

4.2 Equilibrium in the Audit Game

We next state an equilibrium strategy profile for the game $G_{\mathcal{A},k}$. Formally, we present a $(\epsilon_{\mathcal{A},k}, 0)$ -SPE strategy profile, and calculate the value $\epsilon_{\mathcal{A},k}$. The proposed strategy relies on commitment by $\mathcal{D}$ and computation of a single round best response by $\mathcal{A}$. We accordingly refer to this strategy profile as a *simple commitment* strategy profile.

For any equilibrium to be played out with certainty, players must believe that the strategy being used by the other players is the equilibrium strategy. Our proposed strategy profile has features that aim to achieve correct beliefs for the players, even in face of partial rationality. One feature is that $\mathcal{D}$ makes its strategy publicly known, and provides a means to verify that it is playing that strategy. As noted earlier, even though $\mathcal{D}$ acts

after $\mathcal{A}$ does by committing to its strategy with a verification mechanism $\mathcal{D}$ simulates a first move by making the employee believe its commitment with probability one. Thus, we envision the organization making a commitment to stick to its strategy and providing a proof that it follows the strategy. Further, $\mathcal{D}$ making its strategy publicly known follows the general security principle of not making the security mechanisms private [25]. Additionally, the simple commitment strategy profile is an approximate SPE for all values of parameters in any game $G_{\mathcal{A},k}$ and any value of $\mathcal{A}$'s discount factor $\delta_{\mathcal{A}}$. Thus, all employees observe the organization following a consistent strategy further reducing any variability in beliefs about the organization's strategy. Another important feature of the simple commitment strategy profile is the single round best response computation by $\mathcal{A}$ (yielding a single action to play), which is much simpler than optimizing over multiple rounds often yielding many strategies as the solution. Thus, the organization also trusts the employee to make the appropriate decision even if the employee is computationally constrained. The above features of the simple commitment strategy profile makes the strategy simple, which makes it more likely to be followed in the real world.

The main idea behind the definition of our strategy profile is that $\mathcal{D}$ optimizes its utility assuming the best response of $\mathcal{A}$ for a given a^t . That is, $\mathcal{D}$ assumes that $\mathcal{A}$ does not commit any violations when (P, α) is in the deterred region, and systematically commits a violation otherwise (i.e., all of $\mathcal{A}$'s tasks are violations). Further, $\mathcal{D}$ assumes the worst case when the employee (with probability ϵ_{th}) accidentally makes a mistake in the execution of their strategy; in such a case, $\mathcal{D}$ expects all of $\mathcal{A}$'s tasks to be violations, regardless of the values of (P, α) . This is because the distribution D_0^t over violations when $\mathcal{A}$ makes a mistake is unknown. Thus, the expected cost function that $\mathcal{D}$ optimizes (for each total number of tasks a^t) is a linear sum of $(1 - \epsilon_{th})$ times the cost due to best response of $\mathcal{A}$ and ϵ_{th} times the cost when $\mathcal{A}$ commits all violations. The expected cost function is different in the deterred and non-deterred region due to the difference in best response of $\mathcal{A}$ in these two regions. The boundary between the deterred and non-deterred regions is conditioned by the value of the adversary's personal benefit I . We assume that $\mathcal{D}$ learns the value of the personal benefit within an error δI of its actual value, and that $\mathcal{D}$ does not choose actions (P, α) in the region of uncertainty determined by the error δI .

Formally, the expected reward is $E(\mathbf{Rew}_{\mathcal{D}}^t)[0]$ when the adversary commits no violation, and $E(\mathbf{Rew}_{\mathcal{D}}^t)[a^t]$ when all a^t tasks are violations. Both of these expected rewards are functions of P, α ; we do not make that explicit for notational ease. Denote the deterred region determined by the parameter I and the budget $b_{\mathcal{A},k}^t$ as R_D^I and the non-deterred region as R_{ND}^I . Either of these regions may be empty. Denote the region (of uncertainty) between the curves determined by $I + \delta I$ and $I - \delta I$ as $R_{\delta I}^I$. Then the reduced deterred region is given by $R_D^I \setminus R_{\delta I}^I$ and the reduced non-deterred region by $R_{ND}^I \setminus R_{\delta I}^I$. The equilibrium strategy we propose is:

- For each possible number of tasks a^t that can be performed by $\mathcal{A}$, $\mathcal{D}$ constrained by budget $b_{\mathcal{A},k}^t$, assumes the expected utility

$$U_D(P, \alpha) = (1 - \epsilon_{th})E(\mathbf{Rew}_{\mathcal{D}}^t)[0] + \epsilon_{th}E(\mathbf{Rew}_{\mathcal{D}}^t)[a^t] \text{ and} \\ U_{ND}(P, \alpha) = (1 - \epsilon_{th})E(\mathbf{Rew}_{\mathcal{D}}^t)[a^t] + \epsilon_{th}E(\mathbf{Rew}_{\mathcal{D}}^t)[a^t],$$

in $R_D^I \setminus R_{\delta I}^I$ and $R_{ND}^I \setminus R_{\delta I}^I$ respectively. $\mathcal{D}$ calculates the maximum expected utility across the two regions as follows:

$$\begin{aligned} - U_{\max}^D &= \max_{(P, \alpha) \in R_D^I \setminus R_{\delta I}^I} U_D(P, \alpha), U_{\max}^{ND} = \max_{(P, \alpha) \in R_{ND}^I \setminus R_{\delta I}^I} U_{ND}(P, \alpha) \\ - U &= \max(U_{\max}^D, U_{\max}^{ND}) \end{aligned}$$

$\mathcal{D}$ commits to the corresponding maximizer (P, α) for each a^t .

After knowing a^t , $\mathcal{D}$ plays the corresponding (P, α) .

- $\mathcal{A}$ plays her best response (based on the committed action of $\mathcal{D}$), i.e., if she is deterred for all a^t she commits no violations and if she is not deterred for some a^t then all her tasks are violations, and she chooses the a^t that maximizes her utility from violations. But, she also commits mistakes with probability ϵ_{th} , and then the action is determined by distribution D_0^t .

Let $U_{\max}^{D+\delta I} = \max_{(P, \alpha) \in R_D^I \cup R_{\delta I}^I} U_D(P, \alpha)$, $U_{\max}^{ND+\delta I} = \max_{(P, \alpha) \in R_{ND}^I \cup R_{\delta I}^I} U_{ND}(P, \alpha)$, $\delta U^D = U_{\max}^{D+\delta I} - U_{\max}^D$ and $\delta U^{ND} = U_{\max}^{ND+\delta I} - U_{\max}^{ND}$. We have the following result:

Theorem 1. *The simple commitment strategy profile (defined above) is an $(\epsilon_{\mathcal{A},k}, 0)$ -SPE for the game $G_{\mathcal{A},k}$, where $\epsilon_{\mathcal{A},k}$ is*

$$\max \left(\max_{v^t, a^t} (\delta U^D), \max_{v^t, a^t} (\delta U^{ND}) \right) + \epsilon_{th} \max_{\alpha \in [0,1]} \left(\sum_{j=0}^{m-1} \delta_D^j E(r(\vec{O}^t, j)) [U_k, U_k, \alpha] \right)$$

Remark 1. If the value of any parameter of the game (e.g., R_{ext} , R_{int}) is perturbed in a bounded manner, then accounting for that in the analysis yields an $(\epsilon, 0)$ -SPE, but, with ϵ greater than $\epsilon_{\mathcal{A},k}$. This happens because $\mathcal{D}$'s utility is continuous in the parameters.

The proof is in Appendix B. The proof involves showing that the strategy profile has the single stage deviation property. That $\mathcal{A}$ does not profit from deviating is immediate since $\mathcal{A}$ chooses the best response in each round of the game. The bound on profit from deviation for $\mathcal{D}$ has two terms. The first term arises due to $\mathcal{D}$ ignoring the region of uncertainty in maximizing its utility. The maximum difference in utility for the deterred region is $\max_{v^t, a^t} (U_{\max}^{D+\delta I} - U_{\max}^D)$ and for the undeterred region is $\max_{v^t, a^t} (U_{\max}^{ND+\delta I} - U_{\max}^{ND})$. The first term is the maximum of these quantities. The second term arises due to the use of the worst case assumption of all violations out of maximum possible U_k tasks when $\mathcal{A}$ makes a mistake as compared to the case when D_0^t is known. Since $\mathcal{A}$'s choice only affects the violation loss part of $\mathcal{D}$'s utility and mistakes happen with probability ϵ_{th} , the second term is the maximum possible violation loss multiplied by ϵ_{th} .

Numeric applications. The above theorem can be used to calculate concrete values for $\epsilon_{\mathcal{A},k}$ when all parametric functions are instantiated. For example, with the values in Section 3, we obtain $\epsilon_{\mathcal{A},k} = \200 . Assuming $\mathcal{A}$ performs the maximum $U_k = 40$ number of tasks, $\epsilon_{\mathcal{A},k}$ is about 9.5% of the cost of auditing all actions of $\mathcal{A}$ with maximum punishment rate (\$2100), with no violations, and about 3.3% of the cost incurred due to all violations caught externally (\$6000), with no internal auditing or punishment. Similarly, if we assume 70% audit coverage with maximum punishment and four violations, the expected cost for organization is \$2583, which means $\epsilon_{\mathcal{A},k}$ corresponds to about 7.7% of this cost. We present the derivation of value of $\epsilon_{\mathcal{A},k}$ in Claim B in Appendix B. The audit coverage here is for one employee only; hence it can

be as high as 100%. Also, since $\mathcal{G}_{\mathcal{A}}$ is a parallel composition of the games $\mathcal{G}_{\mathcal{A},k}$ for all k , we claim that the simple commitment strategy profile followed for all games $\mathcal{G}_{\mathcal{A},k}$ is a $(\sum_k \epsilon_{\mathcal{A},k}, 0)$ -SPE strategy profile for $G_{\mathcal{A}}$. (See Lemma 1 in Appendix B.1.)

4.3 Comparision with other equilibria

In this section, we compare our proposed strategy with the set of Perfect Public Equilibrium (PPE) strategies. A PPE is the appropriate notion of equilibrium in an imperfect information repeated game with public signals and simultaneous moves. A PPE is quite similar to a SPE; the differences are that histories are sequences of public signals (instead of action profiles) and payoffs are considered in the expected sense. PPE strategy profiles also have the single stage deviation property. As pointed out already, one advantage of the simple commitment strategy is simplicity. As the set of PPE strategies is often infinite, it is difficult for players' beliefs to agree on the strategy being played. However, a commitment by one player to her part of a PPE strategy profile forces that particular PPE to be played. The organization is naturally the player who commits. A *committed utility maximizing* player is one who uses a commitment to force the PPE that yields the maximum payoff to that player. A *privacy preserving* defender is one that chooses a PPE with fewer violations when it has a choice over multiple PPE with the same payoff for the defender. The next theorem shows that simple commitment strategy deters $\mathcal{A}$ as often as the case in which the chosen PPE strategy deters $\mathcal{A}$, assuming the budget allows for deterring the employee and the organization is committed utility maximizing and privacy preserving in choosing PPE equilibrium. Stated succinctly, the simple commitment strategy profile is no worse for privacy protection than choosing the highest utility PPE in scenarios where the organization chooses a PPE strategy that deters the employee.

Theorem 2. *Assume that budget is fixed in every round and is sufficient to deter $\mathcal{A}$, and the number of tasks performed by $\mathcal{A}$ in every round is fixed. Let v_o^* be the maximum PPE payoff that $\mathcal{D}$ can obtain. Further suppose there exists a PPE E_m in which $\mathcal{D}$ always plays some action in the deterred region and the utility for $\mathcal{D}$ with E_m is v_o^* . Then a committed utility maximizing and privacy preserving $\mathcal{D}$ will choose to play E_m . Further, the action in E_m coincides with the action chosen by simple commitment strategy profile in each round.*

5 Budget Allocation

In this section we present two budget allocation mechanisms: one maximizes $\mathcal{D}$'s utility (Section 5.1) and another does the same under accountability constraints (Section 5.2).

5.1 Optimized Budget Allocation

We assume the budget available to $\mathcal{D}$ for all audits is bound by B . Then we must have $\sum_{\mathcal{A},k} \bar{b}_{\mathcal{A}}^t(k) + \text{Cost}(\mathcal{M}) \leq B$, where $\text{Cost}(\mathcal{M})$ is a fixed cost of using the log analysis tool in an audit cycle. Let $B_{\mathcal{M}} = B - \text{Cost}(\mathcal{M})$. Let $\alpha_{\mathcal{A},k}(\vec{b}_{\mathcal{A}}^t(k), \vec{a}_{\mathcal{A}}^t(k))$, $P_{\mathcal{A},k}(\vec{b}_{\mathcal{A}}^t(k), \vec{a}_{\mathcal{A}}^t(k))$

be the equilibrium in game $\mathcal{G}_{\mathcal{A},k}$ for budget $\vec{b}_{\mathcal{A}}^t(k)$ and $\mathcal{A}$'s tasks $\vec{a}_{\mathcal{A}}^t(k)$. Note that we make the dependence on $\vec{b}_{\mathcal{A}}^t(k)$, $\vec{a}_{\mathcal{A}}^t(k)$ explicit here. Let $U(\vec{b}_{\mathcal{A}}^t(k), \vec{a}_{\mathcal{A}}^t(k))$ denote the corresponding expected utility in game $\mathcal{G}_{\mathcal{A},k}$. Observe that in equilibrium, when $\mathcal{A}$ is deterred for all possible $\vec{a}_{\mathcal{A}}^t(k)$ then $\mathcal{A}$ has equal preference for all possible $\vec{a}_{\mathcal{A}}^t(k)$, and otherwise $\mathcal{A}$ chooses the maximum $\vec{a}_{\mathcal{A}}^t(k)$ for which she is undeterred to maximize her utility. Thus, let $BR(\vec{b}_{\mathcal{A}}^t(k))$ be the set of number of tasks all of which are part of best responses of $\mathcal{A}$. Note that the cost functions U_D and U_{ND} in deterred and non-deterred regions are continuous in $\vec{b}_{\mathcal{A}}^t(k)$, since the regions themselves change continuously with change in $\vec{b}_{\mathcal{A}}^t(k)$. Also, by definition they are continuous in $\vec{a}_{\mathcal{A}}^t(k)$. Since U is the maximum of two continuous functions U_D and U_{ND} , using the fact that max of two functions is continuous, we get that U is continuous in both arguments. Then, the optimal allocation of budget is to solve the following non-linear optimization problem

$$\max_{\mathcal{A},k} \sum_{\vec{a}_{\mathcal{A}}^t(k) \in BR(\vec{b}_{\mathcal{A}}^t(k))} \min U(\vec{b}_{\mathcal{A}}^t(k), \vec{a}_{\mathcal{A}}^t(k)) \text{ subject to } \vec{b}_{\mathcal{A}}^t(k) \geq 0 \text{ and } \sum_{\mathcal{A},k} \vec{b}_{\mathcal{A}}^t(k) \leq B_{\mathcal{M}},$$

which maximizes the minimum utility possible over $\mathcal{A}$'s possible best response actions. For example, consider a simple case with two types of tasks: celebrity records accesses and non-celebrity records accesses, and one employee. Assume the utility functions and constants as stated at the end of Section 3, except, it is assumed that it is apriori known that exactly 40 celebrity and 400 non-celebrity accesses would be made and values of some constants (in brackets) are different for celebrity type ($R_{ext} = \$4500, R_{int} = \$300, I = \$6, P_f = 10$) and non-celebrity type ($R_{ext} = \$90, R_{int} = \$30, I = \$0.6, P_f = 5$). Using discrete steps and a brute force search yields a solution of the above optimization problem in which $\mathcal{D}$ would allocate \$1300 to audit celebrity accesses and the remaining \$1200 to audit non-celebrity accesses. As the cost per inspection was assumed \$50 (Section 3), 0.65 fraction of celebrity accesses can be inspected and only 24 out of 400 non-celebrity accesses can be inspected. However, the equilibrium yields that no non-celebrity inspections happen as the employee is non-deterred for the level of non-celebrity inspections possible, and 0.65 fraction of celebrity accesses are inspected.

5.2 Towards Accountable Data Governance

While holding an employee responsible for the violation she causes is natural, it is difficult to define accountability for the organization, as the organization does not commit violations directly. However, the organization influences the actual violator (employee) by the choice of inspections and punishment. We use a simple definition of accountability for the organization, requiring a minimum level of inspection and punishment.

Definition 2. *(($\mathcal{M}, \vec{\alpha}, \vec{P}$)-accountability) An organization satisfies ($\mathcal{M}, \vec{\alpha}, \vec{P}$)-accountability if 1) its log analysis tool $\mathcal{M}'$ satisfies $\mathcal{M}' \geq \mathcal{M}$, 2) its level of inspection satisfies $\vec{\alpha}' \geq \vec{\alpha}$, and 3) its punishment rate satisfies $\vec{P}' \geq \vec{P}$.*

Our definition assumes a partial ordering over log analysis tools $\mathcal{M}$. This partial ordering could be given from empirically computed accuracy μ estimates for each log analysis tool (e.g., we could say that $\mathcal{M}_1 \geq \mathcal{M}_2$ if $\mathcal{M}_1$ is at least as accurate as $\mathcal{M}_2$ for each type of access k). The dependence of accountability on $\mathcal{M}$ is required as a

better performing tool can detect the same expected number of violations as another tool with worse performance, with a lower inspection level α . We envision the above accountability being proven by the organization to a trusted third party external auditor (e.g., Government) by means of a formal proof, in the same manner as commitment is demonstrated to the employee.

To satisfy $(\mathcal{M}, \vec{\alpha}, \vec{P})$ -accountability an organization must add the following constraints to its optimization problem from the last sub-section:

$\min_{\vec{a}_{\mathcal{A}}^t(k) \in BR(\vec{b}_{\mathcal{A}}^t)} \alpha_{\mathcal{A},k}(\vec{b}_{\mathcal{A}}^t(k), \vec{a}_{\mathcal{A}}^t(k)) > \vec{\alpha}(k)$ and $\min_{\vec{a}_{\mathcal{A}}^t(k) \in BR(\vec{b}_{\mathcal{A}}^t)} P_{\mathcal{A},k}(\vec{b}_{\mathcal{A}}^t(k), \vec{a}_{\mathcal{A}}^t(k)) > \vec{P}(k)$ for all $\mathcal{A}, k$. The first constraint ensures that the minimum number of inspections divided by maximum number of tasks is greater than $\vec{\alpha}(k)$, and the second constraint ensures that the minimum punishment level is higher than $\vec{P}(k)$.

Continuing the example from last sub-section if the minimum α and P is specified as 0.1 and 1.0 for both types of accesses, then $\mathcal{D}$ would allocate \$400 to audit celebrity accesses and the remaining \$2100 to audit non-celebrity accesses. Since the cost per inspection was assumed \$50 (Section 3), 0.2 fraction of celebrity accesses can be inspected and 42 out of 400 non-celebrity accesses can be inspected. However, according to the equilibrium 40 non-celebrity inspections happen at punishment level of 2.0 as the employee is already deterred for that level of non-celebrity inspections. In this case, unlike the non-accountable scenario, the values $\vec{\alpha}, \vec{P}$ ensure that the privacy of common person is being protected even when the organization has more economic incentives to audit celebrity accesses more heavily.

6 Predictions and Interventions

In this section, we use our model to predict observed practices in industry and the effectiveness of public policy interventions in encouraging organizations to adopt accountable data governance practices (i.e., conduct more thorough audits) by analyzing the equilibrium audit strategy P, α under varying parameters. The explanation of observed practices provides evidence that our audit model is not far from reality. We use the values of parameters and instantiation of functions given in Section 3 (unless otherwise noted). We assume that the value of personal benefit I is learned exactly and that P and α take discrete values, with the discrete increments being 0.5 and 0.05, respectively. We also assume for sake of exposition that $u_k = U_k$, i.e., the number of tasks is fixed, there is only one type of violation and the budget is sufficient to do all possible inspections.

Average cost R_{ext} and probability p of external detection of violation. We vary R_{ext} from \$5 to \$3900, with R_{int} fixed at \$300. The results are shown in Figure 2. There are two cases shown in the figure: $p = 0.5$ and $p = 0.9$. The figure shows the equilibria P, α chosen for different values of R_{ext} .

Prediction 1: Increasing R_{ext} and p is an effective way to encourage organizations to audit more. In fact, when $p * R_{ext}$ is low X may not audit at all. Thus, X audits to protect itself from greater loss incurred when violations are caught externally. Surprisingly, the hospital may continue to increase inspection levels (incurring higher cost) beyond the minimum level necessary to deter a rational employee. Hospital X does so because the employee is not fully rational: even in the deterred region there is an ϵ_{th} probability of violations occurring.

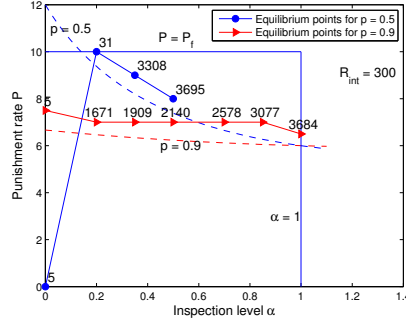


Fig. 2. Separators for two values of external detection probability p indicated by dashed lines. Equilibrium punishment and inspection rates (P, α) marked on solid lines (see legend) as the reputation loss from external detection R_{ext} varies; the R_{ext} values are labeled above the corresponding equilibrium points.

Suggested Intervention 1: Subject organizations to external audits and fines when violations are detected. For example, by awarding contracts for conducting 150 external audits by 2012 [26], HHS is moving in the right direction by effectively increasing p . This intervention is having an impact: the 2011 Ponemon study on patient privacy [27] states—“Concerns about the threat of upcoming HHS HIPAA audits and investigation has affected changes in patient data privacy and security programs, according to 55 percent of respondents.”

Prediction 2: Interventions that increase the expected loss for both external and internal detection of violations are not as effective in increasing auditing as those that increase expected loss for external detection of violations only. Table 2 shows the equilibrium inspection level as R_{ext} and R_{int} are both increased at the same rate. While the inspection level may initially increase, it quickly reaches a peak. As an example, consider the principle of breach detection notification used in many data breach laws [28]. The effect of breach detection notification is to increase both R_{int} and R_{ext} since notification happens for all breaches. While there isn’t sufficient data for our model to predict whether these laws are less effective than external audits (see suggested study below), prior empirical analysis [28] indicate that the benefit in breach detection from these laws is only about 6% (after adjusting for increased reporting of breaches due to the law itself).

Suggested study: An empirical study that separately reports costs incurred when violations are internally detected from those that are externally detected would be useful in quantifying and comparing the effectiveness of interventions. Existing studies either do not speak of these distinct categories of costs [19, 28] or hint at the importance of this distinction without reporting numbers [16, 17].

Punishment loss factor e and personal benefit I . *Prediction 3: Employees with higher value for e (e.g., doctors have higher e ; suspending a doctor is costlier for the hospital than suspending a nurse) will have lower punishment levels.* If punishments were free, i.e., $e = 0$, (an unrealistic assumption) X will always keep the punishment rate at maximum according to our model. At higher punishment rates ($e = 1000$), X will favor increasing inspections rather than increasing the punishment level P (see Table 1 in Appendix A). While we do not know of an industry-wide study on this topic, there is evidence of such phenomena occurring in hospitals. For example, in 2011 Vermont’s Office of Professional Regulation, which licenses nurses, investigated 53 allegations of drug diversion by nurses and disciplined 20. In the same year, the Vermont Board

of Medical Practice, which regulates doctors, listed 11 board actions against licensed physicians for a variety of offenses. However, only one doctor had his license revoked while the rest were allowed to continue practicing [7].

Prediction 4: Employees who cannot be deterred are not punished. When the personal benefit of the employee I is high, our model predicts that X chooses the punishment rate $P = 0$ (because this employee cannot be deterred at all) and increases inspection as R_{ext} increases to minimize the impact of violations by catching them inside (see Table 4 in Appendix A). Note that this is true only for violations that are not very costly (as is the case for our choice of costs). If the expected violation cost is more than the value generated by the employee, then it is better to fire the non-deterred employee (see Appendix B.2).

Audit cost C and performance factor μ of log analysis tool.

Prediction 5: If audit cost C decreases or the performance μ of log analysis increases, then the equilibrium inspection level increases. The data supporting this prediction is presented in Table 3 and 5 in Appendix A. Intuitively, it is expected that if the cost of auditing goes down then organizations would audit more, given their fixed budget allocated for auditing. Similarly, a more efficient mechanized audit tool will enable the organization to increase its audit efficiency with the fixed budget. For example, MedAssets claims that Stanford Hospitals and Clinics saved \$4 million by using automated tools for auditing [29].

7 Related Work

Auditing and Accountability: Prior work studies orthogonal questions of algorithmic detection of policy violations [30–33] and blame assignment [34–37]. Feigenbaum et al. [38] report work in progress on formal definitions of accountability capturing the idea that violators are punished with or without identification and mediation with non-zero probability, and punishments are determined based on an understanding of “typical” utility functions. Operational considerations of how to design an accountability mechanism that effectively manages organizational risk is not central to their work. In other work, auditing is employed to revise access control policies when unintended accesses are detected [39–41]. Another line of work uses logical methods for enforcing a class of policies, which cannot be enforced using preventive access control mechanisms, based on evidence recorded in audit logs [42]. Cheng et al. [43, 44] extend access control to by allowing agents access based on risk estimations. A game-theoretic approach of coupling access control with audits of escalated access requests in the framework of a single-shot game is studied by Zhao et al. [45]. These works are fundamentally different from our approach. We are interested in scenarios where access control is not desirable and audits are used to detect violations. We believe that a repeated game can better model the repeated interactions of auditing.

Risk Management and Data Breaches: Our work is an instance of a risk management technique [12, 13] in the context of auditing and accountability. As far as we know, our technique is the first instance of managing risk in auditing using a repeated game formalism. Risk assessment has been extensively used in many areas [10, 11]; the report by American National Standards Institute [15] provides a risk assessment mechanism

for healthcare. Our model also models data breaches that happen due to insider attacks. Reputation has been used to study insider attacks in non-cooperative repeated games [46]; we differ from that work in that the employer-employee interaction is essentially cooperative. Also, the primary purpose of interaction between employer and employee is to accomplish some task (e.g., provide medical care). Privacy is typically a secondary concern. Our model captures this reality by considering the effect of non-audit interactions in parameters like P_f . There are quite a few empirical studies on data breaches and insider attacks [16, 19, 28] and qualitative models of insider attacks [47]. We use these studies to estimate parameters and evaluate the predictions of our model.

8 Conclusion and Future Work

First, as public policy and industry move towards accountability-based privacy governance, the biggest challenge is how to operationalize requirements such as internal enforcement of policies. We believe that principled audit and punishment schemes like the one presented in this paper can inform practical enforcement regimes. Second, a usual complaint against this kind of risk management approach is that there isn't data to estimate the risk parameters. We provide evidence that a number of parameters in the game model can be estimated from prior empirical studies while recognizing the need for more scientific studies with similar goals, and suggest specific studies that can help estimate other parameters. Third, our model makes an interesting prediction that merits further attention: it suggests that we should design interventions that increase the expected loss from external detection of violations significantly more than the expected loss from internal detection.

While our model captures a number of important economic considerations that influence the design of audit mechanisms, there is much room for further refinement. For example, the model does not handle colluding adversaries nor does it account for detection of violations in audit rounds other than the one in which the violation was committed. Also, our treatment of accountable data governance leaves open questions about the trade-off between utility maximization and privacy protection. Moving forward, we plan to generalize our model, explore the space of policy interventions to encourage accountable data governance, and address normative questions such as what are appropriate levels of inspections and punishments for accountable data governance.

References

1. Center for Information Policy Leadership: Accountability-Based Privacy Governance Project ((accessed May 1,2012))
2. The White House: Consumer Data Privacy in a Networked World: A Framework for Protecting Privacy and Promoting Innovation in the Global Digital Economy ((accessed May 1,2012))
3. Fairwarning: Industry Best Practices for Patient Privacy in Electronic Health Records (April 2011)
4. Hulme, G.: Steady Bleed: State of HealthCare Data Breaches (September 2010) InformationWeek.

5. U.S. Department of Health & Human Services: HIPAA enforcement (accessed May 1,2012)
6. Ornstein, C.: Breaches in privacy cost Kaiser (May 2009) Available online at: <http://articles.latimes.com/2009/may/15/local/me-privacy15>.
7. Ken Picard: Are Drug-Stealing Nurses Punished More Than Doctors? (2012)
8. Blocki, J., Christin, N., Datta, A., Sinha, A.: Regret minimizing audits: A learning-theoretic basis for privacy protection. In: Computer Security Foundations Symposium. (2011) 312–327
9. Fudenberg, D., Tirole, J.: Game Theory. The MIT Press (1991)
10. PricewaterhouseCoopers: A practical guide to risk assessment (December 2008)
11. Karim H. Vellani: Strategic Healthcare Security, Risk Assessments in the Environment of Care (2008) Report for Wisconsin Healthcare Engineering Association.
12. NIST: Guide for Conducting Risk Assessments (September 2011)
13. Pau-Chen Cheng and Pankaj Rohatgi: IT Security as Risk Management: A Reserach Perspective (April 2008) IBM Research Report.
14. Cole Petrochko: DHC: EHR Data Target for Identity Thieves (December 2011)
15. American National Standards Institute(ANSI)/The Santa Fe Group/Internet Security Alliance: The financial impact of breached protected health information (accessed May 1,2012)
16. Verizon: 2012 Data Breach Investigations Report (2012)
17. Ponemon Institute, LLC: Benchmark Study on Patient Privacy and Data Security (November 2010)
18. Ponemon Institute, LLC: 2011 Cost of Data Breach Study: United States (March 2012)
19. Ponemon Institute, LLC: 2010 Annual Study: U.S. Cost of a Data Breach (March 2011)
20. Ichniowski, C., Shaw, K., Prennushi, G.: The Effects of Human Resource Management Practices on Productivity. Technical Report 5333, National Bureau of Economic Research (November 1995)
21. Hanushek, E.A.: Statistical Methods for Social Scientists. New York: Academic Press (1977)
22. Mailath, G.J., Samuelson, L.: Repeated Games and Reputations: Long-Run Relationships. Oxford University Press, USA (2006)
23. Varian, H.: System reliability and free riding. In: Economics of Information Security (Advances in Information Security, Volume 12). (2004) 1–15
24. Grossklags, J., Christin, N., Chuang, J.: Secure or insure? A game-theoretic analysis of information security games. In: World Wide Web Conference (WWW’08). (2008) 209–218
25. Saltzer, J.H., Schroeder, M.D.: The protection of information in computer systems. Proceedings of the IEEE **63**(9) (1975) 1278–1308
26. U.S. Department of Health & Human Services: HIPAA Privacy and Security Audit Program
27. Ponemon Institute, LLC: Second Annual Benchmark Study on Patient Privacy and Data Security (December 2011)
28. Romanosky, S., Hoffman, D., Acquisti, A.: Empirical analysis of data breach litigation. In: International Conference on Information Systems. (2011)
29. MedAssets: MedAssets Case Sudy: Stanford hospital takes charge of its charge capture process, increasing net revenue by 4 million (2011)
30. Barth, A., Datta, A., Mitchell, J.C., Nissenbaum, H.: Privacy and contextual integrity: Framework and applications. In: IEEE Symposium on Security and Privacy. (2006) 184–198
31. Basin, D.A., Klaedtke, F., Müller, S.: Policy monitoring in first-order temporal logic. In: Computer Aided Verification (CAV). (2010) 1–18
32. Garg, D., Jia, L., Datta, A.: Policy auditing over incomplete logs: theory, implementation and applications. In: ACM Computer and Communications Security (CCS). (2011) 151–162
33. Tschantz, M.C., Datta, A., Wing, J.M.: Formalizing and enforcing purpose requirements in privacy policies. In: IEEE Symposium on Security and Privacy. (2012)
34. Backes, M., Datta, A., Derek, A., Mitchell, J.C., Turuani, M.: Compositional analysis of contract-signing protocols. Theor. Comput. Sci. **367**(1-2) (2006) 33–56

35. Barth, A., Datta, A., Mitchell, J.C., Sundaram, S.: Privacy and utility in business processes. In: Computer Security Foundations Symposium (CSF). (2007) 279–294
36. Jagadeesan, R., Jeffrey, A., Pitcher, C., Riely, J.: Towards a theory of accountability and audit. In: European Symposium on Research in Computer Security. (2009) 152–167
37. Küsters, R., Truderung, T., Vogt, A.: Accountability: definition and relationship to verifiability. In: ACM Conference on Computer and Communications Security. (2010) 526–535
38. Feigenbaum, J., Jaggard, A.D., Wright, R.N.: Towards a formal model of accountability. In: Proceedings of the 2011 workshop on New security paradigms workshop. (2011)
39. Bauer, L., Garriss, S., Reiter, M.K.: Detecting and resolving policy misconfigurations in access-control systems. In: Symposium on Access Control Models and Technologies (SACMAT). (2008) 185–194
40. Vaughan, J.A., Jia, L., Mazurak, K., Zdancewic, S.: Evidence-based audit. In: Computer Security Foundations Symposium (CSF). (2008) 177–191
41. Lampson, B.W.: Computer security in the real world. IEEE Computer **37**(6) (2004) 37–46
42. Cederquist, J.G., Corin, R., Dekker, M.A.C., Etalle, S., den Hartog, J.I., Lenzini, G.: Audit-based compliance control. Int. J. Inf. Sec. **6**(2-3) (2007) 133–151
43. Cheng, P.C., Rohatgi, P., Keser, C., Karger, P.A., Wagner, G.M., Reninger, A.S.: Fuzzy Multi-Level Security : An Experiment on Quantified Risk-Adaptive Access Control. In: Proceedings of the IEEE Symposium on Security and Privacy. (2007)
44. Cheng, P.C., Rohatgi, P.: IT Security as Risk Management: A Research Perspective. IBM Research Report **RC24529** (April 2008)
45. Zhao, X., Johnson, M.E.: Access governance: Flexibility with escalation and audit. In: Hawaii International International Conference on Systems Science (HICSS). (2010) 1–13
46. Zhang, N., Yu, W., Fu, X., Das, S.K.: Towards effective defense against insider attacks: The establishment of defender’s reputation. In: IEEE International Conference on Parallel and Distributed Systems. (2008) 501–508
47. Band, S.R., Cappelli, D.M., Fischer, L.F., Moore, A.P., Shaw, E.D., Trzeciak, R.F.: Comparing Insider IT Sabotage and Espionage: A Model-Based Analysis. Technical Report CMU/SEI-2006-TR-026, Carnegie Mellon University (December 2006)

A Experimental Outcomes Supporting Predictions

R_{ext}	P	α
5 to 443	0	0
443 to 3900	6.5	1

Table 1. P, α for $e = 1000$

R_{ext} and R_{int}	P	α
5 to 26	0	0
26 to 3900	10	0.2

Table 2. P, α for constant (0) difference in R_{int}, R_{ext}

C	P	α
10	6.5	1
20	6.5	1
30	7.0	0.85
40	7.5	0.65
50	8.0	0.5
60	9.5	0.25
70	10.0	0.2

Table 3. P, α for varying C

R_{ext}	P	α
5	0	0
670	0	0.1
685	0	0.35
714	0	0.6
748	0	0.85
790	0	1.0

Table 4. P, α for $I = 50$

μ	P	α
1.0	10.0	0.3
1.2	9.5	0.35
1.3	9.5	0.35
1.40	9.0	0.45
1.5	9.0	0.45
1.6	8.5	0.5
1.7	8.5	0.5

Table 5. P, α for varying μ

B Proofs

Reminder of Theorem 1. *The simple commitment strategy profile (defined above) is an $(\epsilon_{\mathcal{A},k}, 0)$ -SPE for the game $G_{\mathcal{A},k}$, where $\epsilon_{\mathcal{A},k}$ is*

$$\max \left(\max_{v^t, a^t} (U_{\max}^{D+\delta I} - U_{\max}^D), \max_{v^t, a^t} (U_{\max}^{ND+\delta I} - U_{\max}^{ND}) \right) + \epsilon_{th} \max_{\alpha \in [0,1]} \left(\sum_{j=0}^{m-1} \delta_{\mathcal{D}}^j E(r(\vec{\mathbf{O}}^t, j)) [U_k, U_k, \alpha] \right)$$

Proof. First the easy case for the employee: the employee always plays a best response. When deterred she is indifferent among any a^t , so choice of a^t does not matter in that case. Thus, there is 0 benefit for the employee by deviating with the history-independent strategy followed. There are two terms in the $\epsilon_{\mathcal{A},k}$ bound for the organization. The first term bounds the profit from deviation due to the fact that the true I is not known. The second term further bounds the profit from deviation due to the fact that the distribution D_t^0 is unknown.

Note that we have lifted the action space of $\mathcal{D}$ to commitment functions. Thus, we need to compare the given commitment with other commitment functions. First, note that if the regions were known properly, and D_0^t known then it is possible to find the commitment that is optimal cost for each fixed value of a^t . Then, it is enough to bound the difference in utility of the the audit commitment function to this optimal commitment function across all values of a^t . We perform the analysis for any fixed a^t , then taking the maximum over all a^t to bound the difference in utility when $\mathcal{D}$ could move to the optimal commitment. We first compare the audit commitment to itself when the true regions are known, then assuming true regions are known we compare the audit commitment to the optimal commitment. Then using triangle inequality we get the required difference for a fixed a^t . Then using the fact that $\max_x f(x) + g(x) \leq \max_x f(x) + \max_x g(x)$ we get the required bound for all a^t .

Suppose the simple commitment strategy profile finds a point in the region $R_D^I \setminus R_{\delta I}^I$. The largest true deterred region can be $R_D^I \cup R_{\delta I}^I$. Thus, $U_{\max}^{D+\delta I} - U_{\max}^D$ represents the maximum profit the organization could have obtained by deviating to another point

using the true deterred region in simple commitment strategy, with some fixed value of v^t and a^t . Then the maximum taken over v^t and a^t gives the maximum possible profit by deviation for the deterred region if the true deterred region were known. Similar argument shows that the maximum profit by deviation for non-deterred region is $\max_{v^t, a^t} (U_{\max}^{ND+\delta I} - U_{\max}^{ND})$. Thus, the absolute maximum profit from deviation for any region is given by the maximum of these two quantities.

Next, assume that the true deterred region R_D is known, so is the non-deterred region R_{ND} . We have already show above that the maximum profit from deviation that the organization would get using simple commitment strategy with R_D instead of $R_D^I \setminus R_{\delta I}^I$, and R_{ND} instead of $R_{ND}^I \setminus R_{\delta I}^I$. Assume that the true regions are known and simple commitment strategy outputs (P, α) to be played by the organization. We use the simplified notation with the game under consideration being $G_{\mathcal{A}, k}$. Denote by $f(P, \alpha)$ the function $E(\mathbf{Rew}_D^t)[0]$, by $g(P, \alpha)$ the function $E(\mathbf{Rew}_D^t)[a^t]$ and by $h(P, \alpha)$ the function $E(\mathbf{Rew}_D^t)[D_0^t]$. The function maximized by (P, α) is

$$U_D(P, \alpha) = (1 - \epsilon_{th})f(P, \alpha) + \epsilon_{th}g(P, \alpha) ,$$

in R_D and is

$$U_{ND}(P, \alpha) = (1 - \epsilon_{th})g(P, \alpha) + \epsilon_{th}g(P, \alpha)$$

in R_{ND} . Suppose D_0^t was known and the point (P', α') is obtained by maximizing

$$U'_D(P, \alpha) = (1 - \epsilon_{th})f(P, \alpha) + \epsilon_{th}h(P, \alpha)$$

in the R_D region and

$$U'_{ND}(P, \alpha) = (1 - \epsilon_{th})g(P, \alpha) + \epsilon_{th}h(P, \alpha)$$

in the R_{ND} region. We emphasize that the function U' is the true expected utility. Consider two different cases

- (P, α) and (P', α') both lie in the same region, say R_D . Then, the maximum benefit to be gained out of deviation is $U'_D(P', \alpha') - U_D(P, \alpha)$, which is

$$(1 - \epsilon_{th})(f(P', \alpha') - f(P, \alpha)) + \epsilon_{th}(h(P', \alpha') - h(P, \alpha))$$

Also, since $U_D(P, \alpha) \geq U_D(P', \alpha')$ we have

$$\epsilon_{th}(g(P, \alpha) - g(P', \alpha')) \geq (1 - \epsilon_{th})(f(P', \alpha') - f(P, \alpha))$$

Thus, the maximum benefit is upper bounded by

$$\epsilon_{th} (g(P, \alpha) - g(P', \alpha') + h(P', \alpha') - h(P, \alpha)) .$$

The upper bound is same for the non-deterred case, since in that case the function $f(.,.)$ is replaced by $g(.,.)$ in both U and U' and the exact same calculation as above yields the same bound.

- (P, α) and (P', α') both lie in different regions, say R_D and R_{ND} respectively. Then, the maximum benefit to be gained out of deviation is $U'_{ND}(P', \alpha') - U'_D(P, \alpha)$, which is

$$(1 - \epsilon_{th})(g(P', \alpha') - f(P, \alpha)) + \epsilon_{th}(h(P', \alpha') - h(P, \alpha)) .$$

Also, since $U_D(P, \alpha) \geq U_{ND}(P', \alpha')$ we have

$$\epsilon_{th}(g(P, \alpha) - g(P', \alpha')) \geq (1 - \epsilon_{th})(g(P', \alpha') - f(P, \alpha)) .$$

Thus, the maximum benefit is upper bounded by

$$\epsilon_{th}(g(P, \alpha) - g(P', \alpha') + h(P', \alpha') - h(P, \alpha)) .$$

Now suppose that (P, α) and (P', α') lie in R_{ND} and R_D respectively. Then, the maximum benefit to be gained out of deviation is $U'_D(P', \alpha') - U'_{ND}(P, \alpha)$, which is

$$(1 - \epsilon_{th})(f(P', \alpha') - g(P, \alpha)) + \epsilon_{th}(h(P', \alpha') - h(P, \alpha)) .$$

Also, since $U_{ND}(P, \alpha) \geq U_D(P', \alpha')$ we have

$$\epsilon_{th}(g(P, \alpha) - g(P', \alpha')) \geq (1 - \epsilon_{th})(f(P', \alpha') - g(P, \alpha)) .$$

Thus, the maximum benefit is upper bounded by

$$\epsilon_{th}(g(P, \alpha) - g(P', \alpha') + h(P', \alpha') - h(P, \alpha)) .$$

The above cases show that the upper bound for profit from deviation in one round is always

$$\epsilon_{th}(g(P, \alpha) - g(P', \alpha') + h(P', \alpha') - h(P, \alpha)) .$$

Using definition of expected rewards we have

$$g(P, \alpha) = -C(\alpha^t a^t) - e(P^t)$$

$$h(P, \alpha) = -C(\alpha^t a^t) - e(P^t) - \sum_{j=0}^{m-1} \delta_o^j E_{D_0^t}(E(r(\mathbf{O}^t, j)))$$

Note that for any P, α

$$h(P, \alpha) - g(P, \alpha) = - \sum_{j=0}^{m-1} \delta_o^j E_{D_0^t}(E(r(\mathbf{O}^t, j))) \leq 0 ,$$

thus, the upper bound above is further bounded by

$$\epsilon_{th}(g(P, \alpha) - h(P, \alpha)) ,$$

which is given by

$$\epsilon_{th} \left(\sum_{j=0}^{m-1} \delta_D^j E_{D_0^t}(E(r(\mathbf{O}^t, j))) \right) .$$

Observe that the above term is maximized over choice of D_0^t when D_0^t places all probability mass on a^t (for any α), i.e., $v^t = a^t$. Also, the expected value of r should be increasing in v^t (since higher v^t means higher detected violations), and $v^t = a^t$ takes a maximum value of U_k for game $G_{\mathcal{A},k}$. Thus, the above term is upper bounded by

$$\epsilon_{th} \max_{\alpha \in [0,1]} \left(\sum_{j=0}^{m-1} \delta_{\mathcal{D}}^j E(r(\mathbf{O}^t, j)) [U_k, U_k, \alpha] \right)$$

Now, add the two bounds to get maximum profit from deviation in one round. Further, using Property 1 and noting that the strategy is history independent for $G_{\mathcal{A},k}$ we now obtain the desired result.

Claim. Assume function instantiations from Section ?? . Thus, given $\nu(\alpha) = \mu\alpha - (\mu - 1)\alpha^2$, we must have $\mu \leq 2$. Further, assuming $C + 2(R_{int} - R_{ext}p) \geq 0$ and $R_{int} \leq R_{ext}p$, the $\epsilon_{\mathcal{A},k}$ from Theorem 1 is given by $\epsilon_{th} U_k \max(R_{int}, R_{ext}p) + \Delta I_{\mathcal{A},k}$, where $\Delta I_{\mathcal{A},k}$ is

$$2\delta I \min \left(\frac{e}{p}, \frac{U_k C}{\mu(1-p)(I - \delta I)} \right)$$

Using values from end of Section 3 we can get $\epsilon_{th} U_k \max(R_{int}, R_{ext}p) = 0.03 * 40 * 150 = 180$, also, the minimum in $\Delta I_{\mathcal{A},k}$ is for $e/p = 20$. Assuming, $\delta I = 0.5$ (remember $i_0 = 1$, assume the learning reduces region of uncertainty by half), we have $\Delta I_{\mathcal{A},k} = 20$. Thus, we get $\epsilon_{\mathcal{A},k} = \200 .

Proof. Remember that Note that for $\nu(\alpha) \leq 1$ to hold, it must be that $\mu\alpha - (\mu - 1)\alpha^2 \leq 1$ for $\alpha \in [0, 1]$. It can be readily verified that this happens only when $\mu \leq 2$. Remember the linear functions assumption means $C(s^t) = C s^t$ and $e(P^t) = e P^t$.

$$\begin{aligned} \max_{\alpha \in [0,1]} \sum_{j=0}^{m-1} \delta_{\mathcal{D}}^j E(r(\mathbf{O}^t, j)) [U_k, U_k, \alpha] = \\ U_k R_{ext}p + U_k \max_{\alpha \in [0,1]} (R_{int} - R_{ext}p) \nu(\alpha) \end{aligned}$$

The relevant part to maximize can be expanded as

$$(R_{int} - R_{ext}p)(\mu\alpha - (\mu - 1)\alpha^2)$$

For $\mu < 2$, $(\mu\alpha - (\mu - 1)\alpha^2)$ increases with $\alpha \in [0, 1]$ (derivative is positive). Thus, if $R_{int} > R_{ext}p$ then $\alpha = 1$ is the maximizer else $\alpha = 0$ is the maximizer. Then, it is not difficult to conclude that the maximum value is $U_k \max(R_{int}, R_{ext}p)$.

Now observe that, since $\mu < 2$, $\mu \geq \mu(\alpha) \geq 1$ The utility function maximized by the organization given the linear function and the example reputation function is (using simple notation)

$$\begin{aligned} -\epsilon_{th} R_{ext}p a^t - e P^t - \alpha^t a^t C - \nu(\alpha^t) a^t \epsilon_{th} (R_{int} - R_{ext}p), \\ -R_{ext}p a^t - e P^t - \alpha^t a^t C - \nu(\alpha^t) a^t (R_{int} - R_{ext}p) \end{aligned}$$

in the reduced deterred region and reduced non-deterred region respectively. Observe that for the non-deterred case, using assumption $C + 2(R_{int} - R_{ext}p) \geq 0$ implies $C + \mu(\alpha)(R_{int} - R_{ext}p) \geq 0$, since $2 > \mu \geq \mu(\alpha) \geq 1$ and all quantities C , R_{int} and $R_{ext}p$ are positive. Thus, the maximizer in non-deterred region is always 0, 0, irrespective of the value of I , hence the difference in costs is zero for the cases when I is known perfectly and when there is an error δI .

Assume $U_{\max}^{D+\delta I}$ occurs for a point P', α' and $U_{\max}^D$ happens for a point P, α , and learned value of personal benefit is I . The interesting case is when $P', \alpha' \neq P, \alpha$ and P, α lies on the curve defined by $I + \delta I$. Then suppose P', α' lie on the curve defined by $I + \delta I - \zeta$ for $2\delta I \geq \zeta \geq 0$. Suppose P', α'' and P'', α' are points on the curve defined by $I + \delta I$ obtained by drawing straight lines from the point P', α' . Thus, $P' \leq P''$ and $\alpha' \leq \alpha''$. Note that since $P'(\nu(\alpha') + p(1 - \nu(\alpha'))) = I + \delta I - \zeta$ and $\nu(\alpha), p \leq 1$, we can claim that $P' \geq I - \delta I$. Then we have

$$\zeta = P''(\nu(\alpha') + p(1 - \nu(\alpha'))) - P'(\nu(\alpha') + p(1 - \nu(\alpha')))$$

or

$$P'' - P' = \frac{\zeta}{\nu(\alpha')(1 - p) + p} \leq \frac{\zeta}{p}$$

Also,

$$\zeta = P'(\nu(\alpha'') + p(1 - \nu(\alpha''))) - P'(\nu(\alpha') + p(1 - \nu(\alpha')))$$

or

$$\nu(\alpha'') - \nu(\alpha') = \frac{\zeta}{(1 - p)P'} \leq \frac{\zeta}{(1 - p)(I - \delta I)}$$

Note that

$$\nu(\alpha'') - \nu(\alpha') = \mu(\alpha'' - \alpha') - (\mu - 1)(\alpha'' - \alpha')((\alpha'' + \alpha'))$$

thus, $\nu(\alpha'') - \nu(\alpha') > \mu(\alpha'' - \alpha')$ and hence

$$(\alpha'' - \alpha') \leq \frac{\zeta}{\mu(1 - p)(I - \delta I)}$$

Also, $U^D(P, \alpha) > U^D(P'', \alpha')$ and $U^D(P, \alpha) > U^D(P', \alpha'')$ and, $U_{\max}^{D+\delta I} - U_{\max}^D = U^{D+\delta I}(P', \alpha') - U^D(P, \alpha)$ means that

$$U_{\max}^{D+\delta I} - U_{\max}^D \leq \min(U^{D+\delta I}(P', \alpha') - U^D(P'', \alpha'), U^{D+\delta I}(P', \alpha') - U^D(P', \alpha''))$$

Also, $U^{D+\delta I}(P', \alpha') - U^D(P'', \alpha')$ is given by

$$-e(P' - P'') \leq \frac{e\zeta}{p}$$

Also, $U^{D+\delta I}(P', \alpha') - U^D(P', \alpha'')$ is given by

$$-a^t(\alpha' - \alpha'')C - a^t(\nu(\alpha') - \nu(\alpha''))\epsilon_{th}(R_{int} - R_{ext}p)$$

which can be simplified to

$$a^t(\alpha'' - \alpha')(C + (\mu - (\mu - 1)(\alpha'' + \alpha'))\epsilon_{th}(R_{int} - R_{ext}p))$$

Using result $1 \leq \mu \leq 2$, we have $2 \geq \mu - (\mu - 1)(\alpha + \alpha') \geq 0$. Using assumption $C + 2(R_{int} - R_{ext}p) \geq 0$ we can say that $(C + (\mu - (\mu - 1)(\alpha + \alpha'))\epsilon_{th}(R_{int} - R_{ext}p)) \geq 0$. Also, since $R_{int} \leq R_{ext}p$, we have $(C + (\mu - (\mu - 1)(\alpha + \alpha'))\epsilon_{th}(R_{int} - R_{ext}p)) \leq C$. Thus, using the inequalities above and $\zeta \leq 2\delta I$, $U_{\max}^{D+\delta I} - U_{\max}^D$ is less than

$$2\delta I \min\left(\frac{e}{p}, \frac{a^t C}{\mu(1-p)(I - \delta I)}\right)$$

which is maximized for $a^t = U_k$.

Reminder of Theorem 2. Assume that budget is fixed in every round and is sufficient to deter $\mathcal{A}$, and the number of tasks performed by $\mathcal{A}$ in every round is fixed. Let v_o^* be the maximum PPE payoff that $\mathcal{D}$ can obtain. Further suppose there exists a PPE E_m in which $\mathcal{D}$ always plays some action in the deterred region and the utility for $\mathcal{D}$ with E_m is v_o^* . Then a committed utility maximizing and privacy preserving $\mathcal{D}$ will choose to play E_m . Further, the action in E_m coincides with the action chosen by simple commitment strategy profile in each round.

Proof. If $\mathcal{D}$ always plays deterred points in E_{max} then the best option $\mathcal{A}$ has is to perform 0 violations (with an ϵ_{th} probability of maximum violations). Let P, α be the point in the deterred region that provides highest payoff to $\mathcal{D}$ (under 0 violations and ϵ_{th} probability of maximum violations). In E_{max} , $\mathcal{D}$ will always play P, α . This because playing any other point in the deterred region results in a profitable single stage deviation by switching to P, α . Since E_{max} is a PPE, there must exist a punishment strategy (low continuation payoff) that makes switching to a non-deterred point non-profitable for $\mathcal{D}$ [22].

Hence the total discounted payoff in E_{max} is $U_D(P, \alpha, 0)$, thus, $U_D(P, \alpha, 0) = v_o^*$. A commitment based utility maximizing $\mathcal{D}$ would prefer a PPE with the highest payoff, and further assuming she chooses the best privacy preserving PPE among those, $\mathcal{D}$ would want the E_{max} PPE to be played. $\mathcal{D}$ will ensure E_{max} is played by committing to its strategy. Thus, for the class of PPE strategies the point P, α will be chosen with the assumptions of the theorem.

Now, we claim that with the simple commitment strategy the same point P, α will be chosen by $\mathcal{D}$. To prove this, suppose on the contrary that a point P', α' is chosen in the non-deterred region. (Note that no other point in deterred region can be chosen - because P, α is the point in the deterred region that provides highest payoff to the $\mathcal{D}$). Let the maximum number of violations be a . Then it must be the case that $U_{ND}(P', \alpha', a) > U_D(P, \alpha, 0)$. Also, $(P', \alpha'), a$ is a Nash equilibrium of the stage game. This is easily proven by observing that a number of violations maximizes $\mathcal{A}$'s utility when P', α' is played by $\mathcal{D}$ and P', α' maximizes $\mathcal{D}$'s utility when a violations are committed by $\mathcal{A}$. A Nash equilibrium played in all rounds is a PPE [22]. Then the PPE that results from playing the above NE in all rounds yields payoff $U_{ND}(P', \alpha', a)$, but, by assumption that $U_D(P, \alpha, 0) = v_o^*$, we obtain that this PPE payoff is greater than v_o^* . This is a contradiction as v_o^* is the maximum payoff that $\mathcal{D}$ can obtain in any PPE. Thus, with our strategy P, α will be chosen by $\mathcal{D}$.

B.1 Repeated Product Game - Definition and Results

If two players play multiple (repeated) *independent* games in parallel then it is possible to consider a composition of these games which is itself a (repeated) game. By *independent* games we mean that these games are played without any influence from the other games in parallel. We define the composition below for a repeated game.

Definition 3. (*Repeated Product Games*) Let the two players play the independent one-shot stage games $G_1, G_2, \dots, G_n$ in parallel in each round of the corresponding n repeated games. A composition of the n stage games is a single-shot game G given by player i 's ($i = 1, 2$) action space $S_i = S_{1i} \times S_{2i} \times \dots \times S_{ni}$, and the payoff function $r_i(s_1, s_2) = \sum_{j=1}^n r_i(s_{j1}, s_{j2})$ where $s_{ji} \in S_{ji}$ and $s_i \in S_i$. A repeated product game is a repeated game with the stage game in every round given by G .

We can extend the above definition to games with imperfect monitoring and public signaling, similar to the manner in which a standard repeated game is extended. Observe that any strategy σ of a repeated product game can be decomposed into strategies $\sigma_1, \dots, \sigma_n$ of the component games, because of the independence assumption of the component games. This decomposition leads to the following useful results summarized in the lemma below:

Lemma 1. Let RG be a repeated product game with the stage game given by G , such that G is a parallel composition of $G_1, G_2, \dots, G_n$ as defined in Definition 3. Consider a strategy σ of RG with the decomposition into strategy σ_i for each component game. Then

- The strategy σ is a SPE iff the strategy σ_i is a SPE for the repeated game with stage game G_i for all i .
- The strategy σ is an $(\sum_{i=1}^n \epsilon_{1i}, \sum_{i=1}^n \epsilon_{2i})$ -SPE if the strategy σ_i is a $(\epsilon_{1i}, \epsilon_{2i})$ -SPE for the repeated game with stage game G_i for all i .

Proof. For the first case assume σ_i is a SPE for the repeated game with stage game G_i for all i . Since σ is given by $\sigma_1, \dots, \sigma_n$ any unilateral deviation from σ results in a unilateral deviation from one or more of $\sigma_1, \dots, \sigma_n$, suppose it is σ_j . By assumption that is not profitable for repeated game given by the stage game G_j . Since the payoff in G is the sum of payoffs in $G_1, \dots, G_n$ and payoffs of games other than j^{th} game remains same, the deviation is not profitable for G also.

The other direction is very similar. Assume σ is a SPE. Since σ is given by $\sigma_1, \dots, \sigma_n$ any unilateral deviation from σ_j results in a unilateral deviation from σ . By assumption that is not profitable for repeated game given by the stage game G . Since the payoff in G is the sum of payoffs in $G_1, \dots, G_n$ and payoffs of games other than j^{th} game remains same, the deviation is not profitable for G_j also.

Next, for the second part since the payoff of G is the sum of payoff's of G_i 's and any any unilateral deviation from σ results in a unilateral deviation from one or more of $\sigma_1, \dots, \sigma_n$, then it is not difficult to check that the profit from deviation will not more than the sum of profit from deviation in each of the repeated games defined by $G_1, \dots, G_n$. Thus, the maximum profit from deviation for player j is $\sum \epsilon_{ji}$.

B.2 Determining P_f and Punishment

In addition to the action dependent utilities above, the players also get an fixed utility in each round of $\mathcal{G}_A$, which is the salary Sal_A for $\mathcal{A}$ and the value created by the employee $g \times Sal_A$ for $\mathcal{D}$. Note that this is the salary and value created for the duration of one audit cycle. Also, note that this fixed utility is not part of any game $\mathcal{G}_{A,k}$. Let R_k be the maximum loss of reputation possible for violation of type k (when all tasks are violations that are externally detected). We assume that the maximum punishment $\vec{P}_{f,A}(k)$ rate for each type k is proportional to R_k . Since the employee can make mistakes, in the worst case he can lose an expected amount of $\epsilon_{th} \sum_k \vec{P}_{f,A}(k) U_k$. This loss must be less than a fixed fraction net of Sal_A , or else the employee is better off quitting and getting better expected payoff in every round in some other job. Thus, we must have $\epsilon_{th} \sum_k \vec{P}_{f,A}(k) U_k = net \cdot Sal_A$, which yields a value $\vec{P}_{f,A}(k) = R_k net \cdot Sal_A / (\epsilon_{th} \sum_k R_k U_k)$. Observe that an employee with higher salary can be punished more. For example, suppose $\mathcal{A}$ does two types k, k' of tasks such that in every week $U_k = 40, U_{k'} = 400$ and $R_{k'} = 0.5 R_k$ and $net = 0.15$ with weekly salary \$500. Then, $\vec{P}_f(k) = 10.4$ and $\vec{P}_f(k') = 5.2$.

Next, consider the case the the employee is non-deterred for violations of type k . Then suppose the expected loss to the organization in every round for such a case is maximum of $U_k L_k$, where L_k is maximum per violation cost (dependent on α) that can be calculated from our model. In such a case if it happens that $U_k L_k \geq (g - 1) Sal_A$ then the organization obtains no benefit from employing $\mathcal{A}$. Thus, in such a case the organization must fire the employee.