



Truthful Mechanisms for Agents that Value Privacy

Citation

Chen, Yiling, Stephen Chong, Ian A. Kash, Tal Moran, and Salil Vadhan. 2013. Truthful Mechanisms for Agents That Value Privacy. In Proceedings of the Fourteenth ACM Conference on Electronic Commerce - EC '13, 215-232. New York: ACM Press.

Published Version

doi:10.1145/2492002.2482549

Permanent link

<http://nrs.harvard.edu/urn-3:HUL.InstRepos:12560997>

Terms of Use

This article was downloaded from Harvard University's DASH repository, and is made available under the terms and conditions applicable to Other Posted Material, as set forth at <http://nrs.harvard.edu/urn-3:HUL.InstRepos:dash.current.terms-of-use#LAA>

Share Your Story

The Harvard community has made this article openly available.
Please share how this access benefits you. [Submit a story](#).

[Accessibility](#)

Truthful Mechanisms for Agents that Value Privacy

YILING CHEN, Harvard University

STEPHEN CHONG, Harvard University

IAN A. KASH, MSR Cambridge

TAL MORAN, Efi Arazi School of Computer Science, IDC Herzliya

SALIL VADHAN, Harvard University

Recent work has constructed economic mechanisms that are both truthful and differentially private. In these mechanisms, privacy is treated separately from the truthfulness; it is not incorporated in players' utility functions (and doing so has been shown to lead to non-truthfulness in some cases). In this work, we propose a new, general way of modelling privacy in players' utility functions. Specifically, we only assume that if an outcome o has the property that any report of player i would have led to o with approximately the same probability, then o has small privacy cost to player i . We give three mechanisms that are truthful with respect to our modelling of privacy: for an election between two candidates, for a discrete version of the facility location problem, and for a general social choice problem with discrete utilities (via a VCG-like mechanism). As the number n of players increases, the social welfare achieved by our mechanisms approaches optimal (as a fraction of n).

Categories and Subject Descriptors: F.0 [Theory of Computation]: General; J.4 [Social and Behavioral Sciences]: Economics

General Terms: Theory, Economics

Additional Key Words and Phrases: Mechanism design, privacy, differential privacy

1. INTRODUCTION

In this paper, we examine the interaction between mechanism design and differential privacy. In particular, we explicitly model privacy in players' utility functions and design truthful mechanisms with respect to it. Our work is motivated by considerations in both fields.

In mechanism design, it has long been recognized that players may not behave as predicted by traditional incentives analysis out of concerns for privacy: in addition to having preferences about the outcome of a mechanism (e.g., who wins an auction, or where a hospital is located), they may also be concerned about what others learn about their private information (e.g., how much they value the auctioned good, or whether they have some medical condition that makes them care more about the hospital's location). The latter concerns are not modelled in most works on mechanism design, and it is natural to try to bring the new models and techniques of differential privacy to bear on them.

Differential privacy [Dwork et al. 2006] is a notion developed to capture privacy when performing statistical analyses of databases. Informally, a randomized algorithm is differentially private if changing a single individual's data does not "substantially" change the output distribution of the algorithm. Thus, differential privacy is not an absolute notion, but rather a quantitative one that needs to be weighed against other objectives. Indeed, differentially private algorithms typically offer a tradeoff between the level of privacy offered to individuals in a database and the accuracy of

Work begun when all the authors were at the Harvard Center for Computation and Society. Supported in part by a gift from Google, Inc. and NSF grants CCF-0915016 and CNS-1237235.

Author's addresses: Y. Chen, S. Chong, and S. Vadhan, Center for Research on Computation and Society and School of Engineering and Applied Sciences, Harvard University, 33 Oxford Street, Cambridge, MA; I. A. Kash, Microsoft Research Cambridge, 7 J J Thomson Ave, Cambridge CB3 0FB, UK; T. Moran, Efi Arazi School of Computer Science, IDC Herzliya. Permission to make digital or hardcopies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies show this notice on the first page or initial screen of a display along with the full citation. Copyrights for components of this work owned by others than ACM must be honored. Abstracting with credits permitted. To copy otherwise, to republish, to post on servers, to redistribute to lists, or to use any component of this work in other works requires prior specific permission and/or a fee. Permissions may be requested from Publications Dept., ACM, Inc., 2 Penn Plaza, Suite 701, New York, NY 10121-0701 USA, fax +1 (212) 869-0481, or permissions@acm.org.
EC'13, June 16–20, 2013, Philadelphia, USA.

Copyright © 2013 ACM 978-1-4503-1962-1/13/06...\$15.00

statistics computed on the database, which we can think of as a “global” objective to be optimized. However, it is also of interest to consider how privacy should be weighed against the objectives of the individuals themselves. Mechanism design provides a natural setting in which to consider such tradeoffs. Attempting to model and reason about privacy in the context of mechanism design seems likely to lead to an improved understanding about the meaning and value of privacy.

1.1. Previous Work

The first work bringing together differential privacy and mechanism design was by McSherry and Talwar [2007]. They showed how to use differential privacy as a *tool* for mechanism design. By definition, differentially private algorithms are insensitive to individuals’ inputs; a change in a single individual’s input has only a small effect on the output distribution of the algorithm. Thus, if a mechanism is differentially private (and players have bounded utility functions), it immediately follows that the mechanism is *approximately truthful*. That is, reporting untruthfully can only provide a small gain in a player’s utility. With this observation, McSherry and Talwar showed how tools from differential privacy allow construction of approximately truthful mechanisms for many problems, including ones where exact truthfulness is impossible.

However, as pointed out by Nissim et al. [2012b], the approximate truthfulness achieved by McSherry and Talwar [2007] may not be a satisfactory solution concept. While differential privacy can guarantee that a player will gain arbitrarily little by lying, it also makes the potential gain from telling the truth equally small. Thus players may choose to lie in order to protect their privacy. Even worse, as shown by an example in Nissim et al. [2012b], in some cases misreporting is a dominant strategy of the game. Thus, it is difficult to predict the outcome and “global” objectives such as social welfare of differentially private mechanisms. Motivated by this, Nissim et al. [2012b] show how to modify some of the mechanisms of McSherry and Talwar [2007] to provide exact truthfulness. To do so, however, they sacrifice differential privacy and/or augment the standard mechanism design framework so that the mechanism can restrict the way in which players can take advantage of the selected outcome.

A recent paper by Xiao [2011, 2013] shows how to remedy this deficiency and construct mechanisms that simultaneously achieve exact truthfulness and differential privacy. Xiao’s paper also points out that even this combination may not be sufficient for getting players that value privacy to report truthfully. Indeed, exact truthfulness only means that a player *weakly* prefers to tell the truth. Lying might not reduce the player’s utility at all (and differential privacy implies that it can only reduce the player’s utility by at most a small amount). On the other hand, differential privacy does not guarantee “perfect” privacy protection, so it is possible that a player’s concern for privacy may still outweigh the small or zero benefit from being truthful.

To address this, Xiao advocated incorporating privacy directly into the players’ utility functions, and seeking mechanisms that are truthful when taking the combined utilities into account. He proposed to measure privacy cost as the mutual information between a player’s type (assumed to come from some prior distribution) and the outcome of the mechanism.¹ Using this measure, he showed that his mechanism does not remain truthful when incorporating privacy into the utility functions, and left as an open problem to construct mechanisms that do. The difficulty is that we can only incentivize truthfulness by giving players an influence on the outcome, but such an influence also leads to privacy costs, which may incentivize lying. Reconciling this tension requires finding models and mechanisms where the privacy costs can be shown to be quantitatively dominated by the outcome-utility benefits of truthfulness.

1.2. Our Contributions

In this paper, we propose a new, more general way of modelling privacy in players’ utility functions. Unlike Xiao’s mutual information measure, our model does not require assuming a prior on players’

¹The mutual information measure is from the original version of Xiao’s paper [Xiao 2011]. Subsequent to our work, he has revised his model to use a different, prior-free measure of privacy [Xiao 2013].

types, and is instead a pointwise model: we simply assume that if an outcome o has the property that any report of player i would have led to o with approximately the same probability, then o has small privacy cost to player i . One motivation for this assumption is that such an outcome o will induce only a small change in a Bayesian adversary’s beliefs about player i (conditioned on the other players’ reports). (This is inspired by a Bayesian interpretation of differential privacy, due to Dwork and McSherry and described in Kasiviswanathan and Smith [2008].) While Xiao’s mutual information measure is not strictly a special case of our model, we show (in the appendix²) that truthfulness with respect to our modelling implies truthfulness with respect to Xiao’s.

With this modelling, we show that it is possible to construct mechanisms where the privacy costs are dominated by the outcome-utility benefits of truthfulness. Specifically, we give three mechanisms that are truthful with respect to our model of privacy: one for an election between two candidates, one for a discrete version of the facility location problem, and one for general social choice problems with discrete utilities (via a VCG-like mechanism). As the number n of players increases, the social welfare achieved by our mechanisms approaches optimal (as a fraction of n).

To illustrate our ideas, consider an election between two candidates, where each player has a noticeable preference for one candidate over the other (but this preference could potentially be outweighed by privacy concerns). To carry out such an election in a differentially private way, we can use a noisy majority vote. That is, we take a majority rule after adding random noise of magnitude $O(1/\epsilon)$, where ϵ is the “differential privacy parameter” that bounds the influence of each player’s vote on the outcome. Now, as we reduce ϵ , we also reduce the privacy costs experienced by each player (according to our model). However, simply bounding the *overall* privacy cost by a function of ϵ does not suffice to argue truthfulness of the election mechanism. Indeed, consider a highly skewed election, where $2/3$ of the voters prefer the first candidate. Then the chance that any single player’s vote affects the outcome is exponentially small in ϵn , where n is the number of voters. If the player’s overall privacy cost can be as large as ϵ , then we would need to set ϵ very small (namely $\tilde{O}(1/n)$) and hence add a very large amount of noise (sacrificing social welfare) in order to achieve truthfulness. However, we argue that the overall privacy cost of being truthful in a skewed election should also be exponentially small. Indeed, for most settings of the random noise, a player cannot affect the outcome by misreporting, so the privacy cost experienced by the player should be the same regardless of whether the player reported truthfully or not. For the remaining settings of the random noise, in which the player can affect the outcome, reporting truthfully provides the player with a noticeable gain in outcome utility, whereas the privacy costs are bounded by a function of ϵ and hence can be made arbitrarily small. Put differently, the *expected* change in privacy cost experienced by a player that changes her report is bounded by ϵ times the probability that the outcome changes (which is always at most $O(\epsilon)$, but can be much smaller, as in the case of a skewed election). This analysis crucially uses the fact that we adopt a *per-outcome* model of privacy, and can be viewed as a generalization of the result of Dwork et al. [2010] showing that every ϵ -differentially private mechanism actually has expected privacy loss $O(\epsilon^2)$.

Our mechanism for discrete facility location is inspired by Xiao’s mechanism [Xiao 2013] (which he shows to be differentially private and truthful if we don’t incorporate privacy into the utility functions), but with some variation in order to be able to apply an analysis like the above. For both the election and facility location mechanisms, an analysis like the one sketched above establishes *universal truthfulness*—truthfulness for every choice of the mechanism’s random coins. For our VCG-like mechanism for general social choice problems, we need to work a bit harder to also ensure that the payments requested do not compromise privacy, and this leads us to only achieve truthfulness in expectation.

Unlike previous works, we do not treat differential privacy as an end in itself but rather as a means to incentivize truthfulness from agents that value privacy. Thus, we do not necessarily need to set the differential privacy parameter ϵ to be very small (corresponding to very high privacy, but a larger price in social welfare); we only need to set it small enough so that the privacy costs are outweighed

²An appendix with omitted proofs and analyses appears in the version of this paper available on the authors websites.

by the agents' preferences for outcomes. Specifically, our analysis shows that as we decrease ϵ , agents' ability to affect the outcome falls, but their expected privacy cost falls even faster. Thus, it is natural to conclude (as we do) that there is some value of ϵ (which may be large if the agents care much more about the outcome than their privacy) at which the privacy cost is small enough relative to the benefit that agents are willing to report truthfully. Moreover, by taking agents' value for privacy into account in the incentives analysis, we can have greater confidence that the agents will actually report truthfully and achieve the approximately optimal social welfare our analysis predicts.

The appropriate setting of ϵ depends on agents' relative valuation of privacy versus the outcomes, and needs to be determined exogenously to apply our mechanisms. Eliciting agents' values for privacy is an orthogonal problem to the one we consider, and is the subject of a different line of work, starting with Ghosh and Roth [2011] (discussed more in the next section). Since we use a dominant-strategy solution concept, it is not disastrous if our assumption about the privacy valuations (and hence our setting of ϵ) is incorrect for some of the agents; the other agents will still have an incentive to report truthfully.

1.3. Other Related Work

Independently of our work, Nissim et al. [2012a] have considered a related way of modelling privacy in players' utilities and constructed truthful mechanisms under their model. They assume that if *all* outcomes o have the property that no player's report affects the probability of o much (i.e., the mechanism is differentially private), then the *overall* privacy cost of the mechanism is small for every player. This is weaker than our assumption, which requires an analogous bound on the privacy cost for each specific outcome o . Indeed, Nissim et al. [2012a] do not consider a per-outcome model of privacy, and thus do not obtain a reduced privacy cost when a player has a very low probability of affecting the outcome (e.g., a highly skewed election). However, they require assumptions that give the mechanism an ability to reward players for truthfulness (through their concept of "agents' reactions," which can be restricted by the mechanism). For example, in the case of an election or poll between two choices (also considered in their paper), they require that a player *directly* benefits from reporting their true choice (e.g., in a poll to determine which of two magazines is more popular, a player will receive a copy of whichever magazine she votes for, providing her with positive utility at no privacy cost), whereas we consider a more standard election where the players only receive utility for their preferred candidate winning (minus any costs due to privacy). In general, we consider the standard mechanism design setting where truthfulness is only rewarded through the public outcome (and possibly payments), and this brings out the central tension discussed earlier: we can only incentivize truthfulness by giving players an influence on the outcome, but such an influence also leads to privacy costs, which may incentivize lying.

Another recent paper that considers a combination of differential privacy and mechanism design is that of Ghosh and Roth [2011]. They consider a setting where each player has some private information and some value for its privacy (measured in a way related to differential privacy). The goal is to design a mechanism for a data analyst to compute a statistic of the players' private information as accurately as possible, by purchasing data from many players and then performing a differentially private computation. In their model, players may lie about their value for privacy, but they cannot provide false data to the analyst. So they design mechanisms that get players to truthfully report their value for privacy. In contrast, we consider settings where players may lie about their data (their private types), but where they have a direct interest in the outcome of the mechanism, which we use to outweigh their value for privacy (so we do not need to explicitly elicit their value for privacy).

Subsequent to our work, Huang and Kannan [2012] examined the properties of the exponential mechanism [McSherry and Talwar 2007], which can be thought of as noisy version of VCG that is slightly different from the one we study. They showed that, with appropriate payments, this mechanism is truthful, individually rational, approximately efficient, and differentially private, but their model does not incorporate privacy costs into players' utility functions.

We remark that there have also been a number of works that consider secure-computation-like notions of privacy for mechanism design problems (see Naor et al. [1999], Dodis et al. [2000], Izmalkov et al. [2005], Parkes et al. [2008], and Feigenbaum et al. [2010] for some examples). In these works, the goal is to ensure that a distributed implementation of a mechanism does not leak much more information than a centralized implementation by a trusted third party. In our setting, we assume we have a trusted third party to implement the mechanism and are concerned with the information leaked by the outcome itself.

2. BACKGROUND ON MECHANISM DESIGN

In this section, we introduce the standard framework of mechanism design to lay the ground for modelling privacy in the context of mechanism design in next section. We use a running example of an election between two candidates. A (deterministic) mechanism is given by the following components:

- A number n of players. These might be voters in an election between two candidates A and B .
- A set Θ of player types. In the election example, we take $\Theta = \{A, B\}$, where $\theta_i \in \Theta$ indicates which of the two candidates is preferred by voter $i \in [n]$.
- A set O of outcomes. In the election example, we take $O = \{A, B\}$, where the outcome indicates which of the two candidates win. (Note that we do not include the tally of the vote as part of the outcome. This turns out to be significant for privacy.)
- Players' action spaces X_i for all $i \in [n]$. In general, a player's action space can be different from his type space. However, in this paper we view the types in Θ to be values that we expect players to know and report. Hence, we require $X_i = \Theta$ for all $i \in [n]$ (i.e., we restrict to direct revelation mechanisms, which is without loss of generality). In the election example, the action of a player is to vote for A or for B .
- An outcome function $M : X_1 \times \dots \times X_n \rightarrow O$ that determines an outcome given players' actions. Since we require $X_i = \Theta$, the outcome function becomes $M : \Theta^n \rightarrow O$. For example, a majority voting mechanism's function maps votes to the candidate who received a majority of votes.
- Player-specific *utility functions* $U_i : \Theta \times O \rightarrow \mathbb{R}$ for $i = 1, \dots, n$, giving the utility of player i as a function of his type and the outcome.

To simplify notation, we use a mechanism's outcome function to represent the mechanism. That is, a mechanism is denoted $M : \Theta^n \rightarrow O$. The goal of mechanism design is then to design a *mechanism* $M : \Theta^n \rightarrow O$ that takes players' (reported) types and selects an outcome so as to maximize some global objective function (e.g. the sum of the players' utilities, known as *social welfare*) even when players may falsely report their type in order to increase their personal utility. The possibility of players' misreporting is typically handled by designing mechanisms that are *incentive-compatible*, i.e., it is in each player's interest to report their type honestly. A strong formulation of incentive compatibility is the notion of *truthfulness* (a.k.a. dominant-strategy incentive compatibility): for all players i , all types $\theta_i \in \Theta$, all alternative reports $\theta'_i \in \Theta$, and all profiles θ_{-i} of the other players' reports³, we have:

$$U_i(\theta_i, M(\theta_i, \theta_{-i})) \geq U_i(\theta'_i, M(\theta'_i, \theta_{-i})). \quad (1)$$

If Inequality (1) holds for player i (but not necessarily all players), we say that the mechanism is *truthful for player i* . Note that we are using θ_{-i} here as both the type and the report of other players. Since truthfulness must hold for all possible reports of other players, it is without loss of generality to assume that other players report their true type. This is in contrast to the notion of a Nash equilibrium which refers to the incentives of player i under the assumption that other players are using equilibrium strategies.

³We adopt the standard game-theory convention that θ_{-i} refers to all components of the vector θ except the one corresponding to player i , and that (θ_i, θ_{-i}) denotes the vector obtained by putting θ_i in the i 'th component and using θ_{-i} for the rest.

In the election example, it is easy to see that standard majority voting is a truthful mechanism. Changing one's vote to a less-preferred candidate can never increase one's utility (it either does not affect the outcome, or does so in a way that results in lower utility).

In this paper, we will allow randomized mechanisms, which we define as $\mathcal{M} : \Theta^n \times \mathcal{R} \rightarrow \mathcal{O}$, where $\mathcal{R}$ is the probability space from which the mechanism makes its random choices (e.g., all possible sequences of coin tosses used by the mechanism). We write $\mathcal{M}(\theta)$ to denote the random variable obtained by sampling r from $\mathcal{R}$ and evaluating $\mathcal{M}(\theta; r)$. This (non-standard) definition of a randomized mechanism is equivalent to the standard one (where the mechanism is a function from reported types to a distribution over outcomes) and makes our analysis clearer.

For randomized mechanisms, one natural generalization of truthfulness is *truthfulness in expectation*: for all players i , all types θ_i , all utility functions U_i , all reports θ'_i , and all profiles θ_{-i} of the other players' reports, we have:

$$\mathbb{E}[U_i(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}))] \geq \mathbb{E}[U_i(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i}))],$$

where the expectation is taken over the random choices of the mechanism.

A stronger notion is that of *universal truthfulness*: for all players i , all types θ_i and utility functions U_i , all alternative reports θ'_i , and all profiles θ_{-i} of the other players' reports, and all $r \in \mathcal{R}$, we have:

$$U_i(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}; r)) \geq U_i(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i}; r)).$$

Thus $\mathcal{M}$ being universally truthful is equivalent to saying that for every $r \in \mathcal{R}$, $\mathcal{M}(\cdot; r)$ is a deterministic truthful mechanism.

3. MODELLING PRIVACY IN MECHANISM DESIGN

The standard framework of mechanism design does not consider a player's value of privacy. In this section, we incorporate privacy into mechanism design and adapt the definitions of truthfulness accordingly. We continue considering the basic mechanism-design setting from Section 2. However, players now care not only about the outcome of the mechanism, but also what that outcome reveals about their private types. Thus, a player's utility becomes

$$U_i = U_i^{out} + U_i^{priv}, \quad (2)$$

where $U_i^{out} : \Theta \times \mathcal{O} \rightarrow \mathbb{R}$ is player i 's utility for the outcome and U_i^{priv} is player i 's utility associated with privacy or information leakage. Before discussing the form of U_i^{priv} (i.e., what are its inputs), we note that in Equation (2), there is already an implicit assumption that privacy can be measured in units that can be linearly traded with other forms of utility. A more general formulation would allow U_i to be an arbitrary monotone function of U_i^{out} and U_i^{priv} , but we make the standard quasi-linearity assumption for simplicity.

Now, we turn to functional form of U_i^{priv} . First, we note that U_i^{priv} should not just be a function of player i 's type and the outcome. What matters is the *functional relationship* between player i 's reported type and the outcome. For example, a voting mechanism that ignores player i 's vote should have zero privacy cost to player i , but one that uses player i 's vote to entirely determine the outcome may have a large privacy cost. So we will allow U_i^{priv} to depend on the mechanism itself, as well as the reports of other players, since these are what determine the functional relationship between player i 's report and the outcome:

$$U_i^{priv} : \Theta \times \mathcal{O} \times \{\mathcal{M} : \Theta^n \times \mathcal{R} \rightarrow \mathcal{O}\} \times \Theta^{n-1} \rightarrow \mathbb{R}. \quad (3)$$

Thus, when the reports of the n players are $\theta' \in \Theta^n$ and the outcome is o , the utility of player i is

$$U_i(\theta_i, o, \mathcal{M}, \theta'_{-i}) = U_i^{out}(\theta_i, o) + U_i^{priv}(\theta_i, o, \mathcal{M}, \theta'_{-i}).$$

In particular, U_i has the same inputs as U_i^{priv} above, including $\mathcal{M}$. Unlike standard mechanism design, we are not given fixed utility functions and then need to design a mechanism with respect to those utility functions. Our choice of mechanism affects the utility functions too!

Note that we do not assume that U_i^{priv} is always negative (in contrast to Xiao [2011]). In some cases, players may prefer for information about them to be kept secret and in other cases they may prefer for it to be leaked (e.g., in case it is flattering). Thus, U_i^{priv} may be better thought of as “informational utility” rather than a “privacy cost”.

It is significant that we do not allow the U_i^{priv} to depend on the *report* or, more generally, the *strategy* of player i . This is again in contrast to Xiao’s modelling of privacy [Xiao 2011]. We will discuss the motivation for our choice in Section 8, and also show (in the appendix) that despite this difference, truthfulness with respect to our modelling implies truthfulness with respect to Xiao’s modelling.

Clearly no mechanism design would be possible if we make no further assumptions about the U_i^{priv} ’s and allow them to be arbitrary, unknown functions (as their behavior could completely cancel the U_i^{out} ’s). Thus, we will make the natural assumption that U_i^{priv} is small if player i ’s report has little influence on the outcome o . More precisely:

ASSUMPTION 3.1 (PRIVACY-VALUE ASSUMPTION). $\forall \theta \in \Theta^n, o \in O, \mathcal{M} :$

$$\left| U_i^{priv}(\theta_i, o, \mathcal{M}, \theta_{-i}) \right| \leq F_i \left(\max_{\theta'_i, \theta''_i \in \Theta} \frac{\Pr[\mathcal{M}(\theta'_i, \theta_{-i}) = o]}{\Pr[\mathcal{M}(\theta''_i, \theta_{-i}) = o]} \right),$$

where $F_i : [1, \infty) \rightarrow [0, \infty]$ is a privacy-bound function with the property that $F_i(x) \rightarrow 0$ as $x \rightarrow 1$, and the probabilities are taken over the random choices of $\mathcal{M}$.

If the mechanism ignores player i ’s report, the right-hand side of (3.1) is $F_i(1)$, which naturally corresponds to a privacy cost of 0. Thus, we are assuming that the privacy costs satisfy a continuity condition as the mechanism’s dependence on player i ’s report decreases. The privacy-bound function F_i could be the same for all players, but we allow it to depend on the player for generality.

Assumption (3.1) is inspired by the notion of *differential privacy*, which is due to Dinur and Nissim [2003], Dwork and Nissim [2004], Blum et al. [2005], and Dwork et al. [2006]. We restate it in our notation:

Definition 3.2. A mechanism $\mathcal{M} : \Theta^n \times \mathcal{R} \rightarrow O$ is ϵ -differentially private iff

$$\forall \theta_{-i} \in \Theta^{n-1}, o \in O \quad \max_{\theta'_i, \theta''_i \in \Theta} \frac{\Pr[\mathcal{M}(\theta'_i, \theta_{-i}) = o]}{\Pr[\mathcal{M}(\theta''_i, \theta_{-i}) = o]} \leq e^\epsilon.$$

By inspection of Assumption (3.1) and Definition 3.2, we have the following result.

PROPOSITION 3.3. If $\mathcal{M}$ is ϵ -differentially private, then for all players i whose utility functions satisfy Assumption (3.1), all $\theta_{-i} \in \Theta^{n-1}$, and $o \in O$, we have $\left| U_i^{priv}(\theta_i, o, \mathcal{M}, \theta_{-i}) \right| \leq F_i(e^\epsilon)$.

In particular, as we take $\epsilon \rightarrow 0$, the privacy cost of any given outcome tends to 0.

Like differential privacy, Assumption (3.1) makes sense only for randomized mechanisms, and only measures the loss in privacy contributed by Player i ’s report when fixing the reports of the other players. In some cases, it may be that the other players’ reports already reveal a lot of information about player i . See Section 8 for further discussion, interpretation, and critiques of our modelling. With this model, the definitions of truthfulness with privacy are direct analogues of the basic definitions given earlier.

Definition 3.4 (Truthfulness with Privacy). Consider a mechanism design problem with n players, type space Θ , and outcome space O . For a player i with utility function $U_i = U_i^{out} + U_i^{priv}$, we say that a randomized mechanism $\mathcal{M} : \Theta^n \times \mathcal{R} \rightarrow O$ is *truthful in expectation for player i* if for all types $\theta_i \in \Theta_i$, all alternative reports $\theta'_i \in \Theta$ for player i , and all possible profiles θ_{-i} of the other players’ reports, we have:

$$\mathbb{E}[U_i(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}, \theta_{-i})] \geq \mathbb{E}[U_i(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i}), \mathcal{M}, \theta_{-i})].$$

We say that $\mathcal{M}$ is *universally truthful for player i* if the inequality further holds for all $r \in \mathcal{R}$:

$$U_i(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}; r), \mathcal{M}, \theta_{-i}) \geq U_i(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i}; r), \mathcal{M}, \theta_{-i}).$$

Unlike in standard settings, $\mathcal{M}$ being universally truthful does *not* mean that the deterministic mechanisms $\mathcal{M}(\cdot; r)$ are truthful. Indeed, even when we fix r , the privacy utility $U_i^{\text{priv}}(\theta, o, \mathcal{M}, \theta_{-i})$ still depends on the original randomized function $\mathcal{M}$, and the privacy properties of $\mathcal{M}$ would be lost if we publicly revealed r . What universal truthfulness means is that player i would still want to report truthfully even if she knew r but it were kept secret from the rest of the world.

4. PRIVATE TWO-CANDIDATE ELECTIONS

Using Proposition 3.3, we will sometimes be able to obtain truthful mechanisms taking privacy into account by applying tools from differential privacy to mechanisms that are already truthful when ignoring privacy. In this section, we give an illustration of this approach in our example of a two-candidate election.

MECHANISM 4.1. *Differentially private election mechanism*

Input: profile $\theta \in \{A, B\}^n$ of votes, privacy parameter $\epsilon > 0$.

- (1) Choose $r \in \mathbb{Z}$ from a discrete Laplace distribution, namely $\Pr[r = k] \propto \exp(-\epsilon|k|)$.
- (2) If $\#\{i : \theta_i = A\} - \#\{i : \theta_i = B\} \geq r$, output A . Otherwise output B .

We show that for sufficiently small ϵ , this mechanism is truthful for players satisfying Assumption 3.1:

THEOREM 4.2. *Mechanism 4.1 is universally truthful for player i provided that, for some function F_i :*

- (1) *Player i 's privacy utility U_i^{priv} satisfies Assumption 3.1 with privacy bound function F_i , and*
- (2) $U_i^{\text{out}}(\theta_i, \theta_i) - U_i^{\text{out}}(\theta_i, -\theta_i) \geq 2F_i(e^\epsilon)$,

Note that Condition 2 holds for sufficiently small $\epsilon > 0$ (since $F_i(x) \rightarrow 0$ as $x \rightarrow 1$). The setting of ϵ needed to achieve truthfulness depends only on how much the players value their preferred candidate (measured by the left-hand side of Condition 2) and how much they value privacy (measured by the right-hand side of Condition 2), and is independent of the number of players n .

PROOF. Fix the actual type $\theta_i \in \{A, B\}$ of player i , a profile θ_{-i} of reports of the other players, and a choice r for $\mathcal{M}$'s randomness. The only alternate report for player i we need to consider is $\theta'_i = -\theta_i$. Let $o = \mathcal{M}(\theta_i, \theta_{-i}; r)$ and $o' = \mathcal{M}(-\theta_i, \theta_{-i}; r)$. We need to show that $U_i(\theta_i, o, \mathcal{M}, \theta_{-i}) \geq U_i(\theta_i, o', \mathcal{M}, \theta_{-i})$, or equivalently

$$U_i^{\text{out}}(\theta_i, o) - U_i^{\text{out}}(\theta_i, o') \geq U_i^{\text{priv}}(\theta_i, o', \mathcal{M}, \theta_{-i}) - U_i^{\text{priv}}(\theta_i, o, \mathcal{M}, \theta_{-i}). \quad (4)$$

We consider two cases:

Case 1: $o = o'$. In this case, Inequality (4) holds because both the left-hand and right-hand sides are zero.

Case 2: $o \neq o'$. This implies that $o = \theta_i$ and $o' = -\theta_i$. (If player i 's report has any effect on the outcome of the differentially private voting mechanism, then it must be that the outcome equals player i 's report.) Thus the left-hand side of Inequality (4) equals $U_i^{\text{out}}(\theta_i, \theta_i) - U_i^{\text{out}}(\theta_i, -\theta_i)$. By Proposition 3.3, the right-hand side of Inequality (4) is at most $2F(e^\epsilon)$. Thus, Inequality (4) holds by hypothesis. $\square$

Of course, truthfulness is not the only property of interest. After all, a mechanism that is simply a constant function is (weakly) truthful. Another property we would like is economic *efficiency*. Typically, this is defined as maximizing social welfare, the sum of players' utilities. Here we consider the sum of *outcome* utilities for simplicity. As is standard, we normalize players' utilities so that all players are counted equally in measuring the social welfare. In our voting example, we wish to

maximize the number of voters' whose preferred candidates win, which is equivalent to normalizing the left-hand side of Condition 2 in Theorem 4.2 to 1. Standard, deterministic majority voting clearly maximizes this measure of social welfare. Our mechanism achieves approximate efficiency (proofs of efficiency results can be found in the appendix):

PROPOSITION 4.3. *For every profile $\theta \in \Theta^n$ of reports, if we select $o \leftarrow \mathcal{M}(\theta)$ using Mechanism 4.1, then:*

- (1) $\Pr [\#\{i : \theta_i = o\} \leq \max_{o' \in \{A, B\}} \#\{i : \theta_i = o'\} - \Delta] < e^{-\epsilon \Delta}.$
- (2) $\mathbb{E} [\#\{i : \theta_i = o\}] > \max_{o' \in \{A, B\}} \#\{i : \theta_i = o'\} - 1/\epsilon.$

Thus, the number of voters whose preferred candidate wins is within $O(1/\epsilon)$ of optimal, in expectation and with high probability. This deviation is independent of n , the number of players. Thus if we take ϵ to be a constant (as suffices for truthfulness) and let $n \rightarrow \infty$, the economic efficiency approaches optimal, when we consider both as fractions of n . This also holds for vanishing $\epsilon = \epsilon(n)$, provided $\epsilon = \omega(1/n)$. Despite the notation, ϵ need not be tiny, which allows a tradeoff between efficiency and privacy.

This analysis considers the social welfare as a sum of outcome utilities (again normalizing so that everyone values their preferred candidate by one unit of utility more than the other candidate). We can consider the effect of privacy utilities on the social welfare too. By Proposition 3.3, the privacy utilities affect the social welfare by at most $\sum_i F_i(e^\epsilon)$, assuming player i satisfies Assumption 3.1 with privacy bound function F_i . If all players satisfy Assumption 3.1 with the same privacy bound function $F_i = F$, then the effect on social welfare is at most $n \cdot F(e^\epsilon)$. By taking $\epsilon \rightarrow 0$ (e.g., $\epsilon = 1/\sqrt{n}$), the privacy utilities contribute a vanishing fraction of n .

Another desirable property is *individual rationality*: players given the additional option of not participating should still prefer to participate and report truthfully. This property follows from the same argument we used to establish universal truthfulness. By dropping out, the only change in outcome that player i can create is to make her less preferred candidate win. Thus, the same argument as in Theorem 4.2 shows that player i prefers truthful participation to dropping out.

PROPOSITION 4.4. *Under the same assumptions as Theorem 4.2, Mechanism 4.1 is individually rational for player i .*

5. TOOLS FOR PROVING TRUTHFULNESS WITH PRIVACY

The analysis of truthfulness in Theorem 4.2 is quite general. It holds for any differentially private mechanism with the property that if a player can actually change the outcome of the mechanism by reporting untruthfully, then it will have a noticeable negative impact on the player's outcome utility. We abstract this property for use in analyzing our other mechanisms.

LEMMA 5.1. *Consider a mechanism design problem with n players, type space Θ , and outcome space O . Let player i have a utility function $U_i = U_i^{\text{out}} + U_i^{\text{priv}}$ satisfying Assumption 3.1 with privacy bound function F_i . Suppose that randomized mechanism $\mathcal{M} : \Theta^n \rightarrow O$ has the following properties:*

- (1) $\mathcal{M}$ is ϵ -differentially private, and
- (2) *For all possible types θ_i , all profiles θ_{-i} of the other players' reports, all random choices r of $\mathcal{M}$, and all alternative reports θ'_i for player i : if $\mathcal{M}(\theta_i, \theta_{-i}; r) \neq \mathcal{M}(\theta'_i, \theta_{-i}; r)$, then $U_i^{\text{out}}(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}; r)) - U_i^{\text{out}}(\theta'_i, \mathcal{M}(\theta'_i, \theta_{-i}; r)) \geq 2F_i(e^\epsilon)$,*

Then $\mathcal{M}$ is universally truthful for player i .

Lemma 5.1 implicitly requires that players not be indifferent between outcomes (unless there is no way for the player to change one outcome to the other). A condition like this is necessary because otherwise players may be able to find reports that have no effect on their outcome utility but improve their privacy utility. However, we show in Section 7 that in some settings with payments, truthfulness can be achieved even with indifference between outcomes since payments can break ties.

It is also illustrative and useful to consider what happens when we take the expectation over the mechanism's coin tosses. We can upper-bound the privacy utility as follows:

LEMMA 5.2. *Consider a mechanism design problem with n players, type space Θ , and outcome space O . Let player i have type $\theta_i \in \Theta_i$ and a utility function $U_i = U_i^{\text{out}} + U_i^{\text{priv}}$ satisfying Assumption 3.1. Suppose that randomized mechanism $\mathcal{M} : \Theta^n \times \mathcal{R} \rightarrow O$ is ϵ -differentially private. Then for all possible profiles θ_{-i} of the other players' reports, all random choices r of $\mathcal{M}$, and all alternative reports θ'_i for player i , we have*

$$\left| \mathbb{E}[U_i^{\text{priv}}(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}, \theta_{-i})] - \mathbb{E}[U_i^{\text{priv}}(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i}), \mathcal{M}, \theta_{-i})] \right| \leq 2F_i(e^\epsilon) \cdot \text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i})),$$

where SD denotes statistical difference.⁴

PROOF. For every two discrete random variables X and Y taking values in a universe $\mathcal{U}$, and every function $f : \mathcal{U} \rightarrow [-1, 1]$, it holds that $|\mathbb{E}[f(X)] - \mathbb{E}[f(Y)]| \leq 2\text{SD}(X, Y)$. (The f that maximizes the left-hand side sets $f(x) = 1$ when $\Pr[X = x] > \Pr[Y = x]$ and sets $f(x) = -1$ otherwise.) Take $\mathcal{U} = O$, $X = \mathcal{M}(\theta_i, \theta_{-i})$, $Y = \mathcal{M}(\theta'_i, \theta_{-i})$, and $f(o) = U_i^{\text{priv}}(\theta_i, o, \mathcal{M}, \theta_{-i})/F_i(e^\epsilon)$. By Proposition 3.3, we have $f(o) \in [-1, 1]$. $\square$

By this lemma, to establish truthfulness in expectation, it suffices to show that the expected gain in outcome utility from reporting θ_i instead of θ'_i grows proportionally with $\text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i}))$. (Specifically, it should be at least the statistical difference times $2F_i(e^\epsilon)$.) In Lemma 5.1, the gain in outcome utility is related to the statistical difference by coupling the random variables $\mathcal{M}(\theta_i, \theta_{-i})$ and $\mathcal{M}(\theta'_i, \theta_{-i})$ according to the random choices r of $\mathcal{M}$. Indeed, $\Pr_r[\mathcal{M}(\theta_i, \theta_{-i}; r) \neq \mathcal{M}(\theta'_i, \theta_{-i}; r)] \geq \text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i}))$. Thus, if the outcome-utility gain from truthfulness is larger than $2F_i(e^\epsilon)$ whenever $\mathcal{M}(\theta_i, \theta_{-i}; r) \neq \mathcal{M}(\theta'_i, \theta_{-i}; r)$, then we have truthfulness in expectation (indeed, even universal truthfulness).

We note that if $\mathcal{M}$ is differentially private, then $\text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i})) \leq e^\epsilon - 1 = O(\epsilon)$, for small ϵ . By Lemma 5.2, the expected difference in privacy utility between any two reports is at most $O(F_i(e^\epsilon) \cdot \epsilon)$. Thus, ϵ -differential privacy helps us twice, once in bounding the pointwise privacy cost (as $2F_i(e^\epsilon)$), and second in bounding the statistical difference between outcomes. On the other hand, for mechanisms satisfying the conditions of Lemma 5.1, the differential privacy only affects the expected outcome utility by a factor related to the statistical difference. This is why, by taking ϵ sufficiently small, we can ensure that the outcome utility of truthfulness dominates the privacy cost.

Lemma 5.2 is related to, indeed inspired by, existing lemmas used to analyze the composition of differentially private mechanisms. These lemmas state that while differential privacy guarantees a worst case bound of ϵ on the “privacy loss” of all possible outputs, this actually implies an *expected* privacy loss of $O(\epsilon^2)$. Such bounds correspond to the special case of Lemma 5.2 when $F_i = \ln$ and we replace the statistical difference with the upper bound $e^\epsilon - 1$. These $O(\epsilon^2)$ bounds on expected privacy loss were proven first in the case of specific mechanisms by Dinur and Nissim [2003] and Dwork and Nissim [2004], and then in the case of arbitrary differentially private mechanisms by Dwork et al. [2010]. In our case, the $O(\epsilon^2)$ bound does not suffice, and we need the stronger bound expressed in terms of the statistical difference. Consider the differentially private election when the vote is highly skewed (e.g., 2/3 vs. 1/3). Then a player has only an exponentially small probability (over the random choice r of the mechanism) of affecting the outcome, and so the expected outcome utility for voting truthfully is exponentially small. On the other hand, by Lemma 5.2, the expected privacy loss is also exponentially small, so we can still have truthfulness.

⁴The *statistical difference* (aka *total variation distance*) between two discrete random variables X and Y taking values in a universe $\mathcal{U}$ is defined to be $\text{SD}(X, Y) = \max_{S \subseteq \mathcal{U}} |\Pr[X \in S] - \Pr[Y \in S]|$.

6. DISCRETE FACILITY LOCATION

In this section, we apply our framework to discrete facility location. Let $\Theta = \{\ell_1 < \ell_2 < \dots < \ell_q\} \subset [0, 1]$ be a finite set of types indicating player's preferred locations for a facility on the unit interval and $O = [0, 1]$. Players prefer to have the facility located as close to them as possible: $U_i^{out}(\theta_i, o) = -|\theta_i - o|$. For example, the mechanism may be selecting a location for a bus stop along a major highway, and the locations $\ell_1, \dots, \ell_q$ might correspond to cities along the highway where potential bus riders live.

Note that the voting game we previously considered can be represented as the special case where $\Theta = \{0, 1\}$. This problem has a well-known truthful and economically efficient mechanism: select the location of the median report. Xiao [2011] gave a private and truthful mechanism for this problem based on taking the median of a perturbed histogram. His analysis only proved that the mechanism satisfies "approximate differential privacy" (often called (ϵ, δ) differential privacy). To use Proposition 3.3, we need the mechanism to satisfy pure ϵ differential privacy (as in Definition 3.2). Here we do that for a variant of Xiao's mechanism.

MECHANISM 6.1. Differentially private discrete facility location mechanism

Input: profile $\theta \in \Theta^n$ of types, privacy parameter $\epsilon > 0$.

- (1) Construct the histogram $h = (h_1, \dots, h_q)$ of reported type frequencies where h_j is the number of reports θ_i of type ℓ_j and $q = |\Theta|$.
- (2) Choose a random (nonnegative, integer) noise vector $r = (r_1, \dots, r_q) \in \mathbb{N}^q$ where the components r_j are chosen independently such that $\Pr[r_j = k]$ is proportional to $\exp(-\epsilon k/2)$.
- (3) Output the type corresponding to median of the perturbed histogram $h + r$. That is, we output $\ell_{\text{Med}(h+r)}$, where for $z \in \mathbb{N}^q$ we define $\text{Med}(z)$ to be the minimum $k \in [q]$ such that $\sum_{j=1}^k z_j \geq \sum_{j=k+1}^q z_j$.

Xiao's mechanism instead chooses the noise components r_j according to a truncated and shifted Laplace distribution. Specifically, $\Pr[r_j = k]$ is proportional to $\exp((\epsilon/2) \cdot |k - t|)$ for $k = 0, \dots, 2t$ and $\Pr[r_j = k] = 0$ for $k > 2t$, where $t = \Theta(\log(1/\delta)/\epsilon)$. This ensures that the noisy histogram $h + r$ is $(\epsilon, q\delta)$ differentially private, and hence the outcome $\ell_{\text{Med}(h+r)}$ is as well. Our proof (in the appendix) directly analyzes the median, without passing through the histogram. This enables us to achieve pure ϵ differential privacy and use a simpler noise distribution. On the other hand, Xiao's analysis is more general, in that it applies to any mechanism that computes its result based on a noisy histogram.

LEMMA 6.2. Mechanism 6.1 is ϵ -differentially private.

We note that the only property the proof uses about the noise distribution is that $\Pr[r = s] = e^{\epsilon/2} \cdot \Pr[r = f_j(s)]$. This property does not hold for Xiao's noise distribution as described, due to it being truncated above at $2t$, but would hold if his noise distribution was truncated only below.

We next show that this mechanism is truthful and individually rational. (The proof can be found in the appendix.)

THEOREM 6.3. Mechanism 6.1 is universally truthful and individually rational for player i provided that, for some function F_i :

- (1) Player i 's privacy utility U_i^{priv} satisfies Assumption 3.1 with privacy bound function F_i , and
- (2) For all $o, o' \in \Theta$ such that $\theta_i < o < o' < \theta_i$ or $o' > o > \theta_i$, we have $U_i^{out}(\theta_i, o) - U_i^{out}(\theta_i, o') \geq 2F_i(e^\epsilon)$.

In particular, if all players share the standard outcome utility function $U_i^{out}(\theta_i, o) = -|\theta_i - o|$ and have the same privacy bound function $F_i = F$, then the mechanism is universally truthful and individually rational provided that

$$\min_{j \neq k} |\ell_j - \ell_k| \geq 2F(e^\epsilon).$$

So, for a fixed set Θ of player types (preferred locations), we can take ϵ to be a small constant and have truthfulness and individual rationality.

PROPOSITION 6.4. *Suppose that every player i has the standard outcome utility function $U_i^{out}(\theta_i, o) = -|\theta_i - o|$. Then for every profile of types $\theta \in \Theta^n$, if we choose $o \leftarrow \mathcal{M}(\theta)$ using Mechanism 6.1, we have*

- (1) $\Pr \left[\sum_i U_i^{out}(\theta_i, o) \leq \max_{o'} \left(\sum_i U_i^{out}(\theta_i, o') \right) - \Delta \right] \leq q \cdot e^{-\epsilon \Delta / q}.$
- (2) $\mathbb{E} \left[\sum_i U_i^{out}(\theta_i, o) \right] \geq \max_{o'} \left(\sum_i U_i^{out}(\theta_i, o') \right) - O(q/\epsilon).$

Thus, the social welfare is within $\Delta = \tilde{O}(q)/\epsilon$ of optimal, both in expectation and with high probability. Like with Proposition 4.3, these bounds are independent of the number n of participants, so we obtain asymptotically optimal social welfare as $n \rightarrow \infty$. Also like the discussion after Proposition 4.3, by taking $\epsilon = \epsilon(n)$ to be such that $\epsilon = o(1)$ and $\epsilon = \omega(1/n)$ (e.g., $\epsilon = 1/\sqrt{n}$), the sum of privacy utilities is a vanishing fraction of n (for participants satisfying Assumption 3.1 with a common privacy bound function F).

7. GENERAL SOCIAL CHOICE PROBLEMS WITH PAYMENTS

In the preceding two sections we have considered social choice problems where a group needs to choose among a (typically small) set of options with mechanisms that do not use money. In this section, we apply our framework social choice problems where payments are possible using an adaptation of the Vickrey-Clarke-Groves (VCG) mechanism. (This is the setting for which the Groves mechanism was originally designed and unlike in auction settings the number of outcomes is independent of the number of players.) In the general case we examine now, we don't assume any structure on the utility functions (other than discreteness), and thus need to use payments to incentivize players to truthfully reveal their preferences.

Specifically, the type $\theta_i \in \Theta$ of a player will specify a utility $U^{out}(\theta_i, o) \in \{0, 1, \dots, M\}$ for each outcome o from a finite set O . This could correspond, for example, to players having values for outcomes expressible in whole dollars with some upper and lower bounds. This assumption ensures a finite set of types Θ and that if a player changes his reported value it must change by some minimum amount (1 with our particular assumption). Note that this formulation still allows players to be indifferent among outcomes. Given our notion of a type, all players share the same outcome utility function $U_i^{out} = U^{out}$. In order to reason about individual rationality, we also assume that the set of types includes a type $\perp$ that corresponds to not participating (i.e., $U_i^{out}(\perp, o) = 0$ for all o and i). For notational convenience, we assume that $O = \{0, 1, \dots, |O| - 1\}$.

Our goal is to choose the outcome o^* that maximizes social welfare (ignoring privacy), i.e., $o^* = \operatorname{argmax}_{o \in O} \sum_i U^{out}(\theta_i, o)$. A standard way to do so is the Groves mechanism, a special case of the more general VCG mechanism. Each player reports his type and then the optimal outcome o^* is chosen based on the reported types. To ensure truthfulness, each player is charged the externality he imposes on others. If $o_{-i} = \operatorname{argmax}_o \sum_{j \neq i} U^{out}(\theta_j, o)$ is the outcome that would have been chosen without i 's input, then player i makes a payment of

$$P_i = \sum_{j \neq i} \left(U^{out}(\theta_j, o_{-i}) - U^{out}(\theta_j, o^*) \right), \quad (5)$$

for a combined utility of $U^{out}(\theta_i, o^*) - P_i$.

In addition to subtracting payments from player i 's utility as above, we also need to consider the effect of payments on privacy. (The modelling in Section 3 did not consider payments.) While it may be reasonable to treat the payments players make as secret, so that making the payment does not reveal information to others, the amount a player is asked to pay reveals information about the reports of *other* players. Therefore, we will require that the mechanism releases some *public* payment information π that enables all players to compute their payments, i.e., the payment P_i of player i should be a function of θ_i , π , and o^* . For example, π could just be the n -tuple $(P_1, \dots, P_n)$,

which corresponds to making all payments public. But in the VCG mechanism it suffices for π to include the value $V_o = \sum_i U^{out}(\theta_i, o)$ for all outcomes $o \in O$, since

$$P_i = (V_{o_{-i}} - U^{out}(\theta_i, o_{-i})) - (V_{o^*} - U^{out}(\theta_i, o^*)) = \max_o ((U^{out}(\theta_i, o^*) - U^{out}(\theta_i, o)) - (V_{o^*} - V_o)),$$

which can be computed using just the V_o 's, o^* , and θ_i . Moreover, we actually only need to release the differences $V_{o^*} - V_o$, and only need to do so for outcomes o such that $V_{o^*} - V_o \leq M$, since only such outcomes have a chance of achieving the above maximum. (Recall that $U^{out}(\theta_i, o) \in \{0, 1, \dots, M\}$.) This observation forms the basis of our mechanism, which we will show to be truthful for players that value privacy (under Assumption 3.1).

Before stating our mechanism, we summarize how we take payments into account in our modelling. Given reports $\theta' \in \Theta^n$ and randomness r , our mechanism $\mathcal{M}(\theta'; r)$ outputs a pair (o^*, π) , where $o^* \in O$ is the selected outcome and π is “payment information”. Each player then should send payment $P_i = P(\theta'_i, o^*, \pi)$ to the mechanism. (The payment function P is something we design together with the mechanism $\mathcal{M}$.) If player i 's true type is θ_i , then her total utility is:

$$U_i(\theta_i, o^*, \pi, \mathcal{M}, \theta') = U^{out}(\theta_i, o^*) - P(\theta'_i, o^*, \pi) + U_i^{priv}(\theta_i, (o^*, \pi), \mathcal{M}, \theta'_{-i}).$$

Note that we measure the privacy of the pair (o^*, π) , since both are released publicly.

To achieve truthfulness for players that value privacy, we will modify the VCG mechanism described above by adding noise to the values V_o . This yields the following mechanism:

MECHANISM 7.1. *Differentially private VCG mechanism*

Input: profile $\theta \in \Theta^n$ of types, privacy parameter $\epsilon > 0$.

- (1) Choose λ_o from a (discrete) Laplace distribution for each outcome o . Specifically, we set $\Pr[\lambda_o = k] \propto \exp(-(\epsilon \cdot |k|)/(M \cdot |O|))$ for every integer $k \in \mathbb{Z}$.
- (2) Calculate values $V_o = \sum_j U^{out}(\theta_j, o) + \lambda_o + o/|O|$ for each outcome o . (Recall that we set $O = \{0, \dots, |O| - 1\}$. The $o/|O|$ term is introduced in order to break ties.)
- (3) Select outcome $o^* = \arg \max_o V_o$.
- (4) Set the payment information $\pi = \{(o, V_{o^*} - V_o) : V_o \geq V_{o^*} - M\}$.
- (5) Output (o^*, π) .

Each player i then sends a payment of

$$P_i = P(\theta_i, o^*, \pi) = \max_o ((U^{out}(\theta_i, o^*) - U^{out}(\theta_i, o)) - (V_{o^*} - V_o)).$$

By standard results on differential privacy, the tuple of noisy values $\{V_o\}$ is ϵ -differentially private. Since the output (o^*, π) is a function of the V_o 's, the output is also differentially private:

LEMMA 7.2. *Mechanism 7.1 is ϵ -differentially private.*

We now prove that the mechanism is truthful in expectation for players that value privacy (satisfying Assumption 3.1). To do this, we use Lemma 5.2, which shows that by taking ϵ sufficiently small, the expected change in privacy utility from misreporting θ'_i instead of θ_i can be made an arbitrarily small fraction of the statistical difference $\text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i}))$. Thus, to show truthfulness in expectation, it suffices to show that the statistical difference is at most a constant factor larger than the expected decrease in utility from misreporting. That is, we want to show:

$$\begin{aligned} \text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i})) &= O\left(\mathbb{E}[U^{out}(\theta_i, \mathcal{M}(\theta_i, \theta_{-i})) - P(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}))] \right. \\ &\quad \left. - \mathbb{E}[U^{out}(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i})) - P(\theta'_i, \mathcal{M}(\theta'_i, \theta_{-i}))]\right). \end{aligned}$$

To bound the statistical difference, we write $\mathcal{M}(\theta; r) = (\mathcal{M}^1(\theta; r), \mathcal{M}^2(\theta; r))$, where $\mathcal{M}^1$ gives the outcome o^* and $\mathcal{M}^2$ gives the payment information π . Then we have:

$$\begin{aligned} \text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i})) &\leq \Pr_r[\mathcal{M}(\theta_i, \theta_{-i}; r) \neq \mathcal{M}(\theta'_i, \theta_{-i}; r)] \\ &\leq \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r)] + \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) \\ &= \mathcal{M}^1(\theta'_i, \theta_{-i}; r) \wedge \mathcal{M}^2(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^2(\theta'_i, \theta_{-i}; r)]. \end{aligned}$$

The next lemma bounds the statistical difference coming from the outcome:

LEMMA 7.3.

$$\begin{aligned} \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r)] &\leq |O| \cdot (\mathbb{E}[U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta_i, \theta_{-i})) - P(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}))] \\ &\quad - \mathbb{E}[U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta'_i, \theta_{-i})) - P(\theta'_i, \mathcal{M}(\theta_i, \theta_{-i}))]). \end{aligned}$$

PROOF. It suffices to show that for every value of r , we have:

$$\begin{aligned} I[\mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r)] &\leq |O| \cdot (U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta_i, \theta_{-i}; r)) - P(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}; r)) \\ &\quad - U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta'_i, \theta_{-i}; r)) - P(\theta'_i, \mathcal{M}(\theta_i, \theta_{-i}; r))), \end{aligned} \quad (6)$$

where $I[X]$ denotes the indicator for the event X . (Then taking expectation over r yields the desired result.)

If $\mathcal{M}^1(\theta_i, \theta_{-i}; r) = \mathcal{M}^1(\theta'_i, \theta_{-i}; r)$, then both the left-hand and right-hand sides are zero. (Recall that the payment made by player i on an outcome o depends only on the reports of the other players and the randomness of the mechanism.)

So consider a value of r such that $\mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r)$ (i.e., where the indicator is 1). We can treat the $\lambda_o + o/|O|$ term added to each V_o as the report of another player to the standard VCG mechanism. We know that

$$U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta_i, \theta_{-i}; r)) - P(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}; r)) - U^{\text{out}}(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i}; r)) - P(\theta'_i, \mathcal{M}(\theta_i, \theta_{-i}; r)) \geq 0$$

because VCG is incentive compatible for players who don't have a privacy utility. Since the mechanism adds an $o/|O|$ term to V_o to avoid ties, the above inequality is strict. Moreover, the left-hand side is at least $1/|O|$, which establishes Inequality (6).

In more detail, let $o^* = \mathcal{M}^1(\theta_i, \theta_{-i}; r)$ and $o' = \mathcal{M}^1(\theta'_i, \theta_{-i}; r)$ for some $o' \neq o^*$. Write $W_o = \sum_{j \neq i} U_j^{\text{out}}(\theta_j, o) + \lambda_o + o/|O|$ for each outcome o (W_o is just V_o excluding the report of player i), and $o_{-i} = \text{argmax}_o W_o$. Since the mechanism chose o^* on report θ_i , we must have

$$W_{o^*} + U^{\text{out}}(\theta_i, o^*) \geq W_{o'} + U^{\text{out}}(\theta_i, o').$$

Since the fractional parts of the two sides are different multiples of $1/|O|$ (namely $o^*/|O|$ and $o'/|O|$), we have:

$$W_{o^*} + U^{\text{out}}(\theta_i, o^*) \geq W_{o'} + U^{\text{out}}(\theta_i, o') + 1/|O|.$$

Thus:

$$\begin{aligned} U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta_i, \theta_{-i}; r)) - P(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}; r)) &= U^{\text{out}}(\theta_i, o^*) - (W_{o_{-i}} - W_{o^*}) \\ &\geq U^{\text{out}}(\theta_i, o') - (W_{o_{-i}} - W_{o'}) + 1/|O| \\ &= U^{\text{out}}(\theta_i, \mathcal{M}(\theta'_i, \theta_{-i}; r)) - P(\theta'_i, \mathcal{M}(\theta_i, \theta_{-i}; r)) + 1/|O|, \end{aligned}$$

establishing Inequality (6). $\square$

Now we need to prove a similar bound for the probability of misreporting only affecting the payment information π . We note that one trivial solution for handling payments is to only collect payments with a very small probability p , but increase the magnitude of the payments by a factor of $1/p$. In order for payments to not contribute more to the statistical difference than the outcome, we

can take p to be the minimum possible nonzero value of the probability that a misreport can change the outcome (i.e., $\Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r)]$). However, this quantity is exponentially small in n . This would make the magnitude of payments exponentially large, which is undesirable. (Our assumption that players are risk neutral seems unreasonable in such a setting.) However, it turns out that we do not actually need to do this; our mechanism already releases payment information with sufficiently low probability. Indeed, we only release payment information relating to an outcome o when V_o is within M of V_{o^*} , and the probability that this occurs cannot be much larger than the probability that the outcome is changed from o^* to o .

LEMMA 7.4.

$$\begin{aligned} \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) = \mathcal{M}^1(\theta'_i, \theta_{-i}; r) \wedge \mathcal{M}^2(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^2(\theta'_i, \theta_{-i}; r)] \\ \leq 2Me^{\epsilon/|O|} \cdot \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r)]. \end{aligned}$$

PROOF. First observe that

$$\begin{aligned} \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) = \mathcal{M}^1(\theta'_i, \theta_{-i}; r) \wedge \mathcal{M}^2(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^2(\theta'_i, \theta_{-i}; r)] \\ \leq \sum_{o_1 \neq o_2} \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) = \mathcal{M}^1(\theta'_i, \theta_{-i}; r) = o_1 \wedge \mathcal{M}^2(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^2(\theta'_i, \theta_{-i}; r) \text{ on } o_2], \end{aligned}$$

by which we mean that either $(o_2, V_{o_1} - V_{o_2})$ is released in one case but not the other or it is released in both cases but with different values.

Fix o_1 and o_2 as above. If $U^{\text{out}}(\theta_i, o_1) - U^{\text{out}}(\theta_i, o_2) = U^{\text{out}}(\theta'_i, o_1) - U^{\text{out}}(\theta'_i, o_2)$, then $\Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) = \mathcal{M}^1(\theta'_i, \theta_{-i}; r) = o_1 \wedge \mathcal{M}^2(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^2(\theta'_i, \theta_{-i}; r) \text{ on } o_2] = 0$ because the difference between V_{o_1} and V_{o_2} is not changed by the misreporting. So assume that $U^{\text{out}}(\theta_i, o_1) - U^{\text{out}}(\theta_i, o_2) \neq U^{\text{out}}(\theta'_i, o_1) - U^{\text{out}}(\theta'_i, o_2)$; these values must differ by at least 1 due to the discreteness assumption. Fix $\lambda_o = k_o$ for $o \neq o_2$. Denote them as a vector $\lambda_{-o_2} = k_{-o_2}$. Consider some value k_{o_2} such that when $\lambda_{o_2} = k_{o_2}$ we have $\mathcal{M}^1(\theta_i, \theta_{-i}; (k_{o_2}, k_{-o_2})) = \mathcal{M}^1(\theta'_i, \theta_{-i}; (k_{o_2}, k_{-o_2})) = o_1$ and $\mathcal{M}^2(\theta_i, \theta_{-i}; (k_{o_2}, k_{-o_2})) \neq \mathcal{M}^2(\theta'_i, \theta_{-i}; (k_{o_2}, k_{-o_2}))$ on o_2 . (If there is no such k_{o_2} then the event has probability 0 for this choice of k_{-o_2} .) Now consider increasing the value of λ_{o_2} . Let $\hat{k}_{o_2}$ be the minimum value such that either $\mathcal{M}^1(\theta_i, \theta_{-i}; (\hat{k}_{o_2}, k_{-o_2})) = o_2$ or $\mathcal{M}^1(\theta'_i, \theta_{-i}; (\hat{k}_{o_2}, k_{-o_2})) = o_2$. At the first such value of $\hat{k}_{o_2}$, only one of these two events will happen because $U^{\text{out}}(\theta_i, o_1) - U^{\text{out}}(\theta_i, o_2)$ and $U^{\text{out}}(\theta'_i, o_1) - U^{\text{out}}(\theta'_i, o_2)$ differ by at least 1. Moreover, we have $\hat{k}_{o_2} \leq k_{o_2} + M$ because with $\lambda_{o_2} = k_{o_2}$ we have $V_{o_1} - V_{o_2} \leq M$ for either report θ_i or θ'_i . Since $\Pr[\lambda_{o_2} = k] \propto \exp(-\epsilon \cdot |k|/(M \cdot |O|))$, we have $\Pr[\lambda_{o_2} = k_{o_2}] \leq \exp(\epsilon/|O|) \cdot \Pr[\lambda_{o_2} = \hat{k}_{o_2}]$. Furthermore, there can be at most M such values of k_{o_2} . Thus,

$$\begin{aligned} \Pr_r[\lambda_{-o_2} = k_{-o_2} \wedge \mathcal{M}^1(\theta_i, \theta_{-i}; r) = \mathcal{M}^1(\theta'_i, \theta_{-i}; r) = o_1 \wedge \mathcal{M}^2(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^2(\theta'_i, \theta_{-i}; r) \text{ on } o_2] \\ \leq Me^{\epsilon/|O|} \Pr_r[\lambda_{-o_2} = k_{-o_2} \wedge \mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r) \\ \wedge \mathcal{M}^1(\theta_i, \theta_{-i}; r) \in \{o_1, o_2\} \wedge \mathcal{M}^1(\theta'_i, \theta_{-i}; r) \in \{o_1, o_2\}] \end{aligned}$$

Summing over all $o_1 \neq o_2$ and k_{-o_2} gives us the lemma. The factor 2 in the lemma statement is due to the fact that

$$\begin{aligned} \sum_{o_1 \neq o_2, k_{o_2}} \Pr_r[\lambda_{-o_2} = k_{-o_2} \wedge \mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r) \\ \wedge \mathcal{M}^1(\theta_i, \theta_{-i}; r) \in \{o_1, o_2\} \wedge \mathcal{M}^1(\theta'_i, \theta_{-i}; r) \in \{o_1, o_2\}] \\ = 2 \Pr_r[\mathcal{M}^1(\theta_i, \theta_{-i}; r) \neq \mathcal{M}^1(\theta'_i, \theta_{-i}; r)]. \end{aligned}$$

□

Combining Lemmas 7.3 and 7.4, we have

$$\begin{aligned} \text{SD}(\mathcal{M}(\theta_i, \theta_{-i}), \mathcal{M}(\theta'_i, \theta_{-i})) &\leq |O| \cdot (1 + 2Me^{\epsilon/|O|}) \cdot (\mathbb{E}[U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta_i, \theta_{-i})) - P(\theta_i, \mathcal{M}(\theta_i, \theta_{-i}))] \\ &\quad - \mathbb{E}[U^{\text{out}}(\theta_i, \mathcal{M}^1(\theta'_i, \theta_{-i})) - P(\theta'_i, \mathcal{M}(\theta_i, \theta_{-i}))]). \end{aligned}$$

Applying Lemma 5.2 gives us our theorem.

THEOREM 7.5. *Mechanism 7.1 is truthful in expectation and individually rational for player i provided that, for some function F_i :*

- (1) *Player i 's privacy utility U_i^{priv} satisfies Assumption 3.1 with privacy bound function F_i , and*
- (2) $2F_i(e^\epsilon) \cdot |O| \cdot (1 + 2Me^{\epsilon/|O|}) \leq 1$.

In particular, if all players have the same privacy bound function $F_i = F$, it suffices to take ϵ to be a sufficiently small constant depending only on M and $|O|$ (and not the number n of players).

Truthfulness in expectation relies on players being risk neutral in terms of their privacy utility so that it is acceptable that with some low probability, the privacy costs are larger than their utility from the outcome. An alternative approach that does not rely on risk neutrality is to switch from the VCG mechanism to the Expected Externality mechanism. This is a variant on VCG that, rather than charging players the actual externality they impose as in Equation (5), charges them their expected externality

$$E_{\theta \sim p} \left[\sum_{j \neq i} U^{\text{out}}(\theta_j, o_{-i}) - U^{\text{out}}(\theta_j, o^*) \right], \quad (7)$$

where p is a prior distribution over Θ^n , o_{-i} is the outcome that maximizes the sum of outcome utilities of players other than i , and o^* is the outcome that maximizes the sum of outcome utilities when i is included. Essentially, i is charged the expected amount he would have to pay under VCG given the prior over types. Since the amount players are charged is independent of the actual reports of others, collecting payments has no privacy implications. (The proof of Lemma 7.3 shows that if we only consider the privacy cost of the outcome, then we have universal truthfulness.) However, the use of a prior means that the truthfulness guarantee only holds in a Bayes-Nash equilibrium. On the other hand, this mechanism does have other nice properties such as being adaptable to guarantee budget balance.

Finally, we show that Mechanism 7.1 approximately preserves VCG's efficiency.

PROPOSITION 7.6. *For every profile of types $\theta \in \Theta^n$, if we choose $o \leftarrow \mathcal{M}(\theta)$ using Mechanism 7.1, then we have:*

- (1) $\Pr \left[\sum_i U_i^{\text{out}}(\theta_i, o) < \max_{o'} \left(\sum_i U_i^{\text{out}}(\theta_i, o') \right) - \Delta \right] \leq 2|O| \cdot e^{-\epsilon\Delta/(2M|O|)},$
- (2) $\mathbb{E} \left[\sum_i U_i^{\text{out}}(\theta_i, o) \right] \geq \max_{o'} \left(\sum_i U_i^{\text{out}}(\theta_i, o') \right) - O(|O|^2 \cdot M/\epsilon).$

Thus, the social welfare is within $\tilde{O}(|O|^2) \cdot M/\epsilon$ of optimal, both in expectation and with high probability. Like with Proposition 4.3, these bounds are independent of the number n of participants, so we obtain asymptotically optimal social welfare as $n \rightarrow \infty$. Also like the discussion after Proposition 4.3, by taking $\epsilon = \epsilon(n)$ to be such that $\epsilon = o(1)$ and $\epsilon = \omega(1/n)$ (e.g., $\epsilon = 1/\sqrt{n}$), the sum of privacy utilities is also a vanishing fraction of n (for participants satisfying Assumption 3.1 with a common privacy bound function F).

8. DISCUSSION

We conclude by discussing a Bayesian interpretation of our privacy model and several of the model's limitations.

Our modelling of privacy in Section 3 is motivated in part by viewing privacy as a concern about *other's beliefs about you*. Fix a randomized mechanism $\mathcal{M} : \Theta^n \times \mathcal{R} \rightarrow O$, a player $i \in [n]$, and

a profile $\theta_{-i} \in \Theta^{n-1}$ of other player's reports. Suppose that an adversary has a prior T_i on the type of player i , as well as a prior S_i on the strategy $\sigma : \Theta \rightarrow \Theta$ played by player i . Then upon seeing an outcome o from the mechanism, the adversary should replace T_i with a posterior T'_i computed according to Bayes' Rule as follows:

$$\Pr[T'_i = \theta_i] = \Pr[T_i = \theta_i | \mathcal{M}(S_i(T_i), \theta_{-i}) = o] = \Pr[T_i = \theta_i] \cdot \frac{\Pr[\mathcal{M}(S_i(T_i), \theta_{-i}) = o | T_i = \theta_i]}{\Pr[\mathcal{M}(S_i(T_i), \theta_{-i}) = o]}.$$

Thus if we set $x = \max_{\theta', \theta'' \in \Theta} (\Pr[\mathcal{M}(\theta', \theta_{-i}) = o] / \Pr[\mathcal{M}(\theta'', \theta_{-i}) = o])$ (the argument of F_i in Assumption 3.1), then we have

$$x^{-1} \cdot \Pr[T_i = \theta_i] \leq \Pr[T'_i = \theta_i] \leq x \cdot \Pr[T_i = \theta_i].$$

So if x is close to 1, then the posterior T'_i is close to the prior T_i , having the same probability mass functions within a factor of x , and consequently having statistical difference at most $x - 1$. Thus, Assumption 3.1 can be justified by asserting that “if an adversary's beliefs about player i do not change much, then it has a minimal impact on player i 's privacy utility.” One way to think of this is that player i has some smooth value function of the adversary's beliefs about her, and her privacy utility is the difference of the value function after and before the Bayesian updating. This reasoning follows the lines of Bayesian interpretations of differential privacy due to Dwork and McSherry, and described in Kasiviswanathan and Smith [2008].

This Bayesian modelling also explains why we do not include the strategy played by i in the privacy utility function U_i^{priv} . How a Bayesian adversary updates its beliefs about player i based on the outcome do not depend on the actual strategy played by i , but rather on the adversary's beliefs about that strategy, denoted by S_i in the above discussion. Given that our mechanisms are truthful, it is most natural to consider S_i as the truthful strategy (i.e., the identity function). If the Bayesian adversary values possessing a correct belief about players, this is analogous to a notion of equilibrium. If we treat the adversary as another player then if the players report truthfully and the adversary assumes the players report truthfully each is responding optimally to the other. However, if player i can successfully convince the adversary that she will follow some other strategy S_i , then this can be implicitly taken into account in U_i^{priv} . (But if player i further deviates from S_i , this should not be taken into account, since the adversary's beliefs will be updated according to S_i .)

Our modelling of privacy in terms of other's beliefs is subject to several (reasonable) critiques:

- Sometimes a small, continuous change in beliefs can result in discrete choices that have a large impact in someone's life. For example, consider a ranking of potential employees to hire, students to admit, or suitors to marry—a small change in beliefs about a candidate may cause them to drop one place in a ranking, and thereby not get hired, admitted, or married. On the other hand, the candidate typically does not know exactly where such a threshold is and so from their perspective the small change in beliefs could be viewed as causing a small change in the probability of rejection.
- Like in differential privacy, we only consider an adversary's beliefs about player i *given the rest of the database*. (This is implicit in us considering a fixed θ_{-i} in Assumption 3.1.) If an adversary believes that player i 's type is correlated with the other players (e.g., given by a joint prior T on Θ^n), then conditioning on $T_{-i} = \theta_{-i}$ may already dramatically change the adversary's beliefs about player i . For example, if the adversary knew that all n voters in a given precinct prefer the same candidate (but don't know which candidate that is), then conditioning on θ_{-i} tells the adversary who player i prefers. We don't measure the (dis)utility for leaking this kind of information. Indeed, the differentially private election mechanism of Theorem 4.2 will leak the preferred candidate in this example (with high probability).
- The word “privacy” is used in many other ways. Instead of being concerned about other's beliefs, one may be concerned about self-representation (e.g., the effect that reporting a given type may have on one's self-image).

Finally, we remark that for some of the mechanism design settings we consider, privacy may not be a major concern on its own. For example, in a standard, deterministic majority-vote election with a large number of voters, each voter’s privacy may intuitively be protected because the other voters provide sufficient noise to obscure her vote. However, just like in our mechanism, the limited influence that voters have on the outcome also means that they gain very little in outcome utility from being truthful, and it may not be clear which of these considerations (privacy or interest in the outcome) should prevail. Our model and analysis can be viewed as a quantitative argument for truthfulness in such real-life election scenarios. We also note that introducing noise artificially, as in differentially private mechanisms, has significant advantages over relying on natural noise, particularly with respect to composition. (For example, if we run a deterministic majority-vote election twice, with a new voter the second time, and the outcome changes, then the new voter’s preference has been revealed despite any uncertainty about the other voter’s votes. Adding independent artificial noise to each tally solves this problem, and ensures differential privacy.)

ACKNOWLEDGMENTS

This work was inspired by discussions under the Harvard Center for Research on Computation and Society’s “Data Marketplace” project. We are grateful to the other participants in those meetings, including Scott Kominers, David Parkes, Felix Fischer, Ariel Procaccia, Aaron Roth, Latanya Sweeney, and Jon Ullman. We also thank Moshe Babaioff, Kobbi Nissim, and Dave Xiao for helpful discussions and comments.

REFERENCES

- BLUM, A., DWORK, C., MCSHERRY, F., AND NISSIM, K. 2005. Practical privacy: the sulq framework. In *PODS*, C. Li, Ed. ACM, 128–138.
- COVER, T. M. AND THOMAS, J. A. 1991. *Elements of Information Theory* 2nd Ed. Wiley Series in Telecommunications. John Wiley & Sons, Inc.
- DINUR, I. AND NISSIM, K. 2003. Revealing information while preserving privacy. In *PODS*. ACM, 202–210.
- DODIS, Y., HALEVI, S., AND RABIN, T. 2000. A cryptographic solution to a game theoretic problem. In *CRYPTO*, M. Bellare, Ed. LNCS Series, vol. 1880. Springer, 112–130.
- DWORK, C., MCSHERRY, F., NISSIM, K., AND SMITH, A. 2006. Calibrating noise to sensitivity in private data analysis. In *TCC*, S. Halevi and T. Rabin, Eds. LNCS Series, vol. 3876. Springer, 265–284.
- DWORK, C. AND NISSIM, K. 2004. Privacy-preserving datamining on vertically partitioned databases. In *CRYPTO*, M. K. Franklin, Ed. LNCS Series, vol. 3152. Springer, 528–544.
- DWORK, C., ROTHBLUM, G. N., AND VADHAN, S. P. 2010. Boosting and differential privacy. In *FOCS*. IEEE Computer Society, 51–60.
- FEIGENBAUM, J., JAGGARD, A. D., AND SCHAPIRA, M. 2010. Approximate privacy: foundations and quantification (extended abstract). In *ACM EC*. ACM, 167–178.
- GHOSH, A. AND ROTH, A. 2011. Selling privacy at auction. In *ACM EC*. EC ’11. ACM, New York, NY, USA, 199–208.
- HUANG, Z. AND KANNAN, S. 2012. The exponential mechanism for social welfare: Private, truthful, and nearly optimal. In *FOCS*.
- IZMALKOV, S., MICALI, S., AND LEPINSKI, M. 2005. Rational secure computation and ideal mechanism design. In *FOCS*. IEEE Computer Society, 585–595.
- KASIVISWANATHAN, S. P. AND SMITH, A. 2008. A note on differential privacy: Defining resistance to arbitrary side information. *CoRR abs/0803.3946*.
- MCSHERRY, F. AND TALWAR, K. 2007. Mechanism design via differential privacy. In *FOCS*. IEEE Computer Society, 94–103.
- NAOR, M., PINKAS, B., AND SUMNER, R. 1999. Privacy preserving auctions and mechanism design. In *ACM EC*. 129–139.
- NISSIM, K., ORLANDI, C., AND SMORODINSKY, R. 2012a. Privacy-aware mechanism design. In *ACM EC*. 774–789.
- NISSIM, K., SMORODINSKY, R., AND TENNENHOLTZ, M. 2012b. Approximately optimal mechanism design via differential privacy. In *ITCS*, S. Goldwasser, Ed. ACM, 203–213.
- PARKES, D. C., RABIN, M. O., SHIEBER, S. M., AND THORPE, C. 2008. Practical secrecy-preserving, verifiably correct and trustworthy auctions. *Elec. Commerce Res. and Applications* 7, 3, 294–312.
- XIAO, D. 2011. Is privacy compatible with truthfulness? Tech. Rep. 2011/005, Cryptology ePrint Archive.
- XIAO, D. 2013. Is privacy compatible with truthfulness? In *ITCS*, R. D. Kleinberg, Ed. ACM, 67–86.