

Privacy, Prediction, and Allocation

Ben Jacobsen   

Department of Computer Sciences, University of Wisconsin-Madison, WI, USA

Nitin Kohli  

Center for Effective Global Action, University of California Berkeley, CA, USA

Abstract

Algorithmic predictions are increasingly used to inform the allocation of scarce resources. The promise of these methods is that, through machine learning, they can better identify the people who would benefit most from interventions. Recently, however, several works have called this assumption into question by demonstrating the existence of settings where simple, unit-level allocation strategies can meet or even exceed the performance of those based on individual-level targeting. Separately, other works have objected to individual-level targeting on privacy grounds, leading to an unusual situation where a single solution, unit-level targeting, is recommended for reasons of both privacy and utility. Motivated by the desire to fully understand the interplay of privacy and targeting levels, we initiate the study of aid allocation systems that satisfy *differential privacy*, synthesizing existing works on private optimization with the economic models of aid allocation used in the non-private literature. To this end, we investigate private variants of both individual and unit-level allocation strategies in both stochastic and distribution-free settings under a range of constraints on data availability. Through this analysis, we provide clean, interpretable bounds characterizing the tradeoffs between privacy, efficiency, and targeting precision in allocation.

2012 ACM Subject Classification Security and privacy → Human and societal aspects of security and privacy

Keywords and phrases Differential privacy, fair allocation, limits of prediction

Digital Object Identifier 10.4230/LIPIcs.FORC.2026.20

Related Version *Full Version*: <https://arxiv.org/abs/2604.15596>

Funding *Nitin Kohli*: Supported by a grant from the Gates Foundation.

1 Introduction

Public institutions, humanitarian organizations, and private firms increasingly use data to design and deliver interventions across populations. Notable examples can be found across such varied domains as education [12, 52], humanitarian response [3, 5, 40, 41], housing assistance [47], and financial inclusion [10, 11, 41]. A common theme across these applications is the growing use of personal and behavioral data to decide how resources and benefits should be allocated to improve welfare outcomes.

When designing such interventions, the level at which aid is targeted has far-reaching policy implications [1, 65]. Two choices are typical [58]: individual-level allocation (ILA) draws on detailed data about specific people – either directly observed or predicted – to assign benefits one person at a time. It is often praised for precision and efficiency [39] but criticized for its reliance on sensitive personal data [55]. Unit-level allocation (ULA) instead uses aggregated data about groups, such as geographic areas or demographic clusters, to allocate resources collectively. It requires less personal information and may offer greater inclusivity, but this comes at the cost of coarser allocation precision [4].

Privacy concerns arise at all levels of this process [18, 45, 63], but the focus of this work is on the privacy of the allocation itself. For example, when allocations are based on sensitive characteristics such as poverty status, the mere fact that an individual receives an anti-poverty benefit discloses that characteristic, which has been shown to provoke shame



© Ben Jacobsen and Nitin Kohli;
licensed under Creative Commons License CC-BY 4.0
7th Symposium on Foundations of Responsible Computing (FORC 2026).
Editor: Huijia (Rachel) Lin; Article No. 20; pp. 20:1–20:24



Leibniz International Proceedings in Informatics
LIPIC Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

or jealousy among peers and to expose individuals to additional, well-documented privacy risks [35]. A promising approach to address this problem is to use *differential privacy* (DP), a rigorous mathematical framework for limiting the information that an algorithm leaks about its input. But although a number of prior works have investigated private variants of ILA from an optimization perspective [29, 30, 17], and there are examples of systems that have employed DP for both ULA [40] and ILA [41] in specific applications, there is limited understanding on the fundamental tradeoffs between privacy and efficiency in allocation systems more broadly. In particular, it is not clear to what extent the addition of privacy requirements should influence how administrators navigate the fundamental tension between spending their limited budget on collecting more data vs. helping more people, which has been much discussed in recent works [58, 53, 23].

In this paper, we help to close this gap by systematically analyzing the tradeoffs between ILA and ULA across different levels of data availability, subject to DP constraints. As a warm-up, we first consider an idealized setting with full access to each individual’s welfare where the only challenge is to satisfy DP (Section 3). We then turn our attention to more challenging settings where administrators have limited access to individual welfare scores and must balance paying for greater data access with allocating greater amounts of aid. In Section 4, we consider an administrator that can pay to sample welfare scores directly, and design strategies with robust worst-case bounds on regret even when those scores are fixed arbitrarily. Then, in Section 5, we introduce distributional assumptions by considering the setting (increasingly common in practice [3, 5, 41]) where an administrator has access to auxiliary information (features) on the whole population and samples individual welfare scores to train a predictive model. Building on sample complexity bounds for DP-SCO [21], we investigate when and how approaches based on statistical modeling can improve the efficiency of private allocations.

For each setting we consider, we derive clean, interpretable bounds describing the asymptotic regimes where one strategy outperforms the other (Sections 4.3 and 5.3). Perhaps surprisingly, we find that the choice of DP parameter has little impact on the optimal choice of strategy: we present simple, practical DP algorithms for ULA and ILA whose regret relative to non-private baselines vanishes asymptotically. Beyond the level of privacy, we find that ULA is more efficient than ILA when budgets are small, data is expensive, average welfare is low, inequality is high, and individual welfare is difficult to predict. These results are consistent with and extend those of other recent works [58, 52, 53], indicating that many of the defining features of ILA and ULA carry over when privacy constraints are introduced. Taken together, our results provide the first general framework for navigating tradeoffs between privacy, efficiency, and targeting precision in the design of systems for fair allocation.

2 Preliminaries

2.1 Differential Privacy

Differential privacy (DP) is a mathematical notion of privacy designed for statistical tasks [20]. Informally speaking, a randomized mechanism satisfies DP if changing one individual’s data doesn’t change the distribution of the mechanism’s outputs by too much. Our privacy guarantees are stated in terms of zCDP, but can be easily translated to Rényi DP or approximate DP using standard results [15].

► **Definition 1** (ψ -zero-Concentrated Differential Privacy [15]). *A mechanism $M : \mathcal{X}^n \rightarrow \mathcal{Y}$ is ψ -zCDP if, for all $x, x' \in \mathcal{X}^n$ differing on a single entry and all $\alpha \in (1, \infty)$, we have $D_\alpha(M(x) \| M(x')) \leq \psi\alpha$, where $D_\alpha(M(x) \| M(x'))$ is the α -Rényi divergence between the distribution of $M(x)$ and the distribution of $M(x')$.*

Definition 1 turns out to be slightly too strict for computing allocations [29]: if the output of our algorithm is a full allocation, then satisfying DP implies that the aid allocated to an individual cannot depend too much on that same individual's welfare, which is clearly incompatible with effective targeting. It is therefore standard to use a slight relaxation of DP called Joint DP [37]. Essentially, instead of publishing our full allocation, we separately tell each individual whether or not they were chosen to receive aid. Satisfying Joint DP then requires that changing one individual's data won't change the joint distribution of *everybody else's* outputs by too much, but the original individual's output can change arbitrarily.

► **Definition 2** (Joint ψ -zero-Concentrated Differential Privacy; Adapted from [37]). *Let $M : \mathcal{X}^n \rightarrow \mathcal{Y}^n$ be a randomized mechanism that produces n outputs $y = (y_1, \dots, y_n)$ when provided $x = (x_1, \dots, x_n)$. Let $M(x)_{-i}$ denote the output of $M(x)$ without the i^{th} component. We say that M satisfies Joint ψ -zCDP if for all $i \in [n]$, all $x, x' \in \mathcal{X}^n$ differing in their i^{th} entry, and all $\alpha \in (1, \infty)$, we have $D_\alpha(M(x)_{-i} \| M(x')_{-i}) \leq \psi\alpha$.*

A classic result from Joint DP is the *Billboard Lemma*, which allows us to reduce the task of constructing Joint DP mechanisms to standard DP mechanisms.

► **Lemma 3** (Billboard Lemma; Adapted from [29] Lemma 1). *Let $M : \mathcal{X}^n \rightarrow \mathcal{Y}$ satisfy ψ -zCDP. Then for any function $f : \mathcal{Y} \times \mathcal{X} \rightarrow \mathcal{Z}$, the mechanism $M' : \mathcal{X}^n \rightarrow \mathcal{Z}^n$ defined by $M'(x)_i = f(M(x), x_i)$ satisfies Joint ψ -zCDP.*

2.2 Related Work

2.2.1 Allocation Based on Prediction

There is now a substantial body of work analyzing the promises and pitfalls of predictive systems for resource allocation. Some papers emphasize the benefits of predictive models in improving the efficiency and targeting of limited resources [39, 5, 12, 4, 47]. Others have studied the theoretical tradeoffs of prediction-guided allocation [16, 33, 58, 53] and empirically investigated the value of prediction in allocation across different contexts [23, 52, 22].

A distinct line of works has criticized the concept of predictive optimization on normative grounds, arguing contra [39] that such systems often fail on their own terms [62, 9, 43]. In parallel, a growing body of research has suggested that many aspects of life trajectories may be fundamentally difficult to predict regardless of the method used [56, 44, 38].

Most relevant to our study is the framework proposed by Shirali et al. [58], who characterize the conditions necessary for ILA to outperform ULA. They develop a simple mathematical model for aid allocation and find that ILA only outperforms ULA when inequality across units is low and the budget is high. We extend their framework through a detailed examination of the impact of data constraints on efficiency, including constraints motivated by privacy.

2.2.2 Differentially Private Allocation and Matching

Joint DP (Definition 2) was first formalized by [37] in the context of mechanism design for mediated games, but it was quickly observed that the same definition could also be used to compute private matchings and allocations. The pioneering work in this direction is [29], which studied a setting where n agents each have private valuations over (subsets of) k goods. Subsequent works have extended these results to a wide range of matching and convex optimization problems under a variety of assumptions [30, 17, 19]. The main feature that distinguishes these works from our own is that they all begin with the assumption that agents' valuations are known in advance, whereas we study tradeoffs between different strategies

for learning those valuations. Our work can therefore be seen as a bridge between existing theoretical works on Joint DP allocation and the growing non-private literature, described above, that examines the tradeoffs inherent to targeted aid allocation.

2.3 Our Basic Model

Our work expands on the model of allocation presented in Shirali et al [58]. Our main goal in this section is to introduce the key terms and concepts of this model; the motivation and limitations of our modeling choices are discussed in greater detail in Section 6.

To start out, we assume that our population is divided M disjoint units $U_1, \dots, U_M$, each with population N , giving a total population size of $P = MN$. Each individual $i \in [P]$ has a welfare $w_i \in [0, 1]$, and our goal is to maximize the sum of all welfare scores. To this end, we can pay a fixed cost of c to allocate aid to an individual, subject to a budget constraint of B .

We use τ_i to denote the *treatment effect* of aiding individual i , defined as the marginal increase in their welfare after our intervention. Following [58], we assume that $\tau_i = \min(1, w_i + \delta_w) - w_i$ for some fixed δ_w , i.e. treatment effects are homogeneous and welfare is capped at 1. We say that an individual i is “low-welfare” if $w_i \leq 1 - \delta_w$ and “high-welfare” otherwise.

Deciding on an allocation can be seen as a matching problem between an administrator with $k = B/c$ indivisible, exchangeable goods and P potential recipients, where the value of allocating aid to individual i is τ_i . We represent allocations with index sets, and define the *value* of an allocation with index set $\mathcal{I}$ given a welfare vector w as $\text{VAL}(\mathcal{I}, w) := \sum_{i \in \mathcal{I}} \tau_i$. This quantity represents the total improvement in welfare that occurs because of our intervention. To evaluate a particular allocation, we define its *regret* by how far it falls short of the best value we could have achieved:

$$\text{REGRET}(\mathcal{I}, w; k) = \max_{|\mathcal{J}|=k} \text{VAL}(\mathcal{J}, w) - \text{VAL}(\mathcal{I}, w) \quad (1)$$

Abusing notation, we will also define the regret of an algorithm A that outputs an allocation by $\text{REGRET}(A, w; k) := \mathbb{E}_{\mathcal{I} \sim A(w)} [\text{REGRET}(\mathcal{I}, w; k)]$, and likewise for $\text{VAL}(A, w)$. The *normalized regret* is defined as $\overline{\text{REGRET}}(\cdot, w; k) := \frac{1}{\delta_w} \text{REGRET}(\cdot, w; k)$. In some contexts, we will consider welfare scores drawn i.i.d. from some distribution $\mathcal{D}$, in which case the expectation is taken over the randomness of both our algorithm and the input.

When discussing specific algorithms, we annotate a strategy with a hat when it is derived from estimated welfare scores, as in $\widehat{\text{ILA}}$, and with a tilde when it satisfies DP, as in $\widetilde{\text{ILA}}$. The exact nature of the estimation or privacy protection should always be clear from context.

2.3.1 Individual-Level Allocation (ILA)

ILA aids individuals in ascending order by welfare. Let $s(i)$ denote the index of the individual with the i th lowest welfare. Then the value of the optimal non-private allocation is given by $\text{VAL}(\text{ILA}, w) := \sum_{i=1}^k \tau_{s(i)}$. This allocation is exactly optimal and so has zero regret by definition, but is also clearly not privacy-preserving.

To remedy this, we design a variant of ILA that satisfies Joint DP (Definition 2). We begin by sorting individuals in ascending order by either true welfare scores (if they are known) or by an estimate of the probability $\eta_i := \mathbb{P}[w_i > 1 - \delta_w]$ (if w must be estimated). We then learn a global decision rule $f : [0, 1] \rightarrow \{0, 1\}$ subject to central DP and output a single bit $f(w_i)$ (resp. $f(\hat{\eta}_i)$) to each individual indicating whether they will receive aid.

To match the non-private baseline, we want f to be a good approximation of the non-private decision rule $f(w) \approx \mathbb{I}[w \leq w_{s(k)}]$, taking care to handle situations where scores may be highly concentrated or tied. The main technical challenge here is to ensure that our global

budget constraint is still satisfied with high probability, given that each individual allocation decision must be made locally. When analyzing the regret of private variants of ILA, we will make frequent use of the following lemma:

► **Lemma 4.** *Let $\hat{\eta}_i$ denote an estimate for the probability that $w_i \leq 1 - \delta_w$, η_i be the true probability, and let $\hat{s}(j)$ be the index of the person with the j th lowest estimated probability. Suppose that our allocation only aids $k' \leq k$ individuals based on ascending estimated probabilities of having low welfare. Then we have:*

$$\overline{\text{REGRET}}(\widehat{\text{ILA}}, w; k) \leq (k - k') + k' \bar{\eta}_{k'} + \sum_{i \in \mathcal{I}} |\hat{\eta}_i - \eta_i| \quad (2)$$

where $\mathcal{I} := \{i = s(j) = \hat{s}(j') \mid (j \leq k' < j') \vee (j' \leq k' < j)\}$ is the set of individuals wrongly included or excluded from our allocation because of estimation error, and $\bar{\eta}_{k'} := \frac{1}{k'} \sum_{i=1}^{k'} \eta_{s(i)}$ is the probability that a randomly selected individual among those with the k' lowest conditional probabilities of having high welfare does, in fact, have high welfare.

► **Remark 5.** In Section 3 and Section 4, we treat welfare scores as fixed, in which case $\eta_i \in \{0, 1\}$ and $\bar{\eta}_{k'}$ is typically equal to 0. We consider probabilistic welfare in Section 5.

2.3.2 Unit-Level Allocation (ULA)

ULA targets aid allocation to the worst-off units without distinguishing between individuals within the same unit. Specifically, let $\rho_j = \frac{1}{|U_j|} \sum_{i \in U_j} \mathbb{I}[w_i > 1 - \delta_w]$ denote the probability that a randomly selected individual in unit j has high welfare. With no privacy constraints, ULA greedily allocates aid to everybody in the units with the lowest ρ_j scores until its budget is exhausted. To avoid complications from rounding, we permit ULA to allocate resources uniformly at random within the last chosen unit whenever k is not evenly divisible by N . We denote the total number of units treated by $K = \lceil k/N \rceil$.

Letting T_j denote the average treatment effect within unit j and using $s(j)$ to denote the index of the unit with the j th lowest value in ρ , the value of the non-private allocation is:

$$\text{VAL}(\text{ULA}, w) := \sum_{j=1}^{\lfloor k/N \rfloor} N T_{s(j)} + (k \bmod N) T_{s(\lfloor k/N \rfloor + 1)} \geq k \left(\frac{1}{K} \sum_{j=1}^K T_{s(j)} \right) \quad (3)$$

In Lemma 4.1 of [58], the authors show that $\sum_{j=1}^K T_{s(j)} \geq \sum_{j=1}^K \delta_w (1 - \rho_{s(j)})$. This then implies that $\overline{\text{REGRET}}(\text{ULA}, w; k) \leq k \bar{\rho}_K$, where $\bar{\rho}_K := \frac{1}{K} \sum_{j=1}^K \rho_{s(j)}$.

The private variant of ULA operates in the same way, except that it first computes $\tilde{\rho}$, a ψ -zCDP estimate of vector of unit profiles ρ , and then allocates aid in ascending order based on this estimate. The following lemma allows us to decompose the worst-case regret of private ULA into the worst-case regret of non-private ULA plus the error from privacy:

► **Lemma 6.** *Let x_j^* denote the number of individuals in unit j that receive aid under $\text{ULA}(w)$, and let $\tilde{x}_j^*$ denote the same for $\widetilde{\text{ULA}}(w)$. Then for any $p, q \in [1, \infty]$ such that $1/p + 1/q = 1$, we can upper-bound the worst-case normalized regret of the allocation based on $\tilde{\rho}$ by:*

$$\overline{\text{REGRET}}(\widetilde{\text{ULA}}, w; k) \leq k \bar{\rho}_K + \|x^* - \tilde{x}^*\|_p \|\tilde{\rho} - \rho\|_q \quad (4)$$

In particular, we can choose between any of the following bounds for this privacy term:

$$\begin{cases} 2 \min(k, P - k) \max_j |\tilde{\rho}_j - \rho_j| & p = 1, q = \infty \\ \sqrt{2 \min(k, P - k) N \sum_{j=1}^M (\tilde{\rho}_j - \rho_j)^2} & p = 2, q = 2 \\ N \sum_{j=1}^M |\tilde{\rho}_j - \rho_j| & p = \infty, q = 1 \end{cases} \quad (5)$$

► **Remark 7.** Lemma 6 generalizes the connection between $\text{VAL}(\text{ULA}, w)$ and error in unit-level statistics presented in [58, Appendix D.1], which corresponds to the case of $p = \infty, q = 1$.

2.3.3 Random Allocation

Finally, we will also frequently consider the baseline strategy of allocating aid to individuals sampled uniformly without replacement from the population (RAND), which may be preferable to either ILA or ULA when privacy or budgetary constraints are too strict to permit any form of effective targeting.

► **Lemma 8.** Define $\bar{\rho}_M = \frac{1}{M} \sum_{j=1}^M \rho_j$, and let X be a hypergeometric random variable with population size P , number of samples k , and number of successes $P(1 - \bar{\rho}_M)$. Then $\overline{\text{REGRET}}(\text{RAND}, w; k) \leq \mathbb{E}[k - X] = k\bar{\rho}_M$

2.3.4 The Role of Inequality

Prior works [58, 52] have identified inequality between units as a major factor determining the (in)efficiency of ULA. One way to understand the connection is to think of inequality between units as a measure of how *informative* unit membership is; inequality is high exactly when unit membership is a useful proxy for welfare. The following lemma formalizes this intuition by quantifying inequality in terms of the *Gini coefficient* of the unit profile vector, which was also studied in [58]:

► **Lemma 9.** Let $\rho \in [0, 1]^M$ be a unit profile vector with Gini coefficient G , defined as $G = \frac{1}{2M^2\bar{\rho}_M} \sum_{i,j} |\rho_i - \rho_j| = \frac{1}{M^2\bar{\rho}_M} \sum_{1 \leq i \leq j \leq M} (\rho_{s(j)} - \rho_{s(i)})$. Then $\bar{\rho}_M - \bar{\rho}_K \geq \frac{GM\bar{\rho}_M}{M-1} > G\bar{\rho}_M$

Recall that the worst-case regret of ULA scales linearly with $\bar{\rho}_K$, while the worst-case regret of random allocation scales with $\bar{\rho}_M$. Lemma 9 can therefore be interpreted as either a lower bound on the marginal value of ULA relative to RAND, or, by rearranging, as an upper-bound on the marginal regret of ULA relative to ILA.

3 Private Allocation when Welfare is Directly Observable

As a warm-up, we assume that the true welfare scores of all individuals are already known to the administrator. In this simple setting, the only challenge is to ensure that the final allocation we select satisfies Joint DP without exceeding our budget constraints.

3.1 ILA

Our strategy for privatizing ILA is presented in Algorithm 1; we compute a private estimate of the empirical CDF of welfare scores, from which we derive a high-probability lower bound on $w_{s(k)}$, the maximum welfare treated by the optimal allocation. Our own allocation assigns aid to everybody whose (possibly noisy) welfare falls below this bound. The following results show that Algorithm 1 is private and provide high probability upper and lower bounds on the number of individuals treated:

► **Lemma 10.** Algorithm 1 satisfies Joint ψ -zCDP for any valid w, ψ, θ, s , and k .

► **Lemma 11** (Algorithm 1 does not exceed budget whp). Fix w, ψ, θ, s , and k , and let ℓ_{j^*} be the threshold returned by Algorithm 1. Then with probability at least $1 - \beta/2$, $\sum_{i=1}^P \mathbb{I}[\hat{w}_i \leq \ell_{j^*}] \leq k$.

■ **Algorithm 1** Joint ψ -zCDP ILA.

Input:

- Vector of welfare scores $w = [w_1, \dots, w_n] \subseteq [0, 1]^n$
- Privacy parameter $\psi > 0$; bin width $\theta > 0$; noise parameter $s \geq 0$; aid budget k ; failure probability β

- 1 Sample noisy welfare scores $\hat{w}_i \sim \text{UNIF}[w_i - s, w_i + s]$ for $1 \leq i \leq n$
- 2 Initialize interval $\mathcal{R} = [-s, 1 + s]$ and define $B_1 = [\ell_1, r_1] = [-s, -s + \theta]$.
- 3 Recursively define bins $B_j = (\ell_j, r_j] = (r_{j-1}, r_{j-1} + \theta]$ for $2 \leq j \leq (1 + 2s)/\theta$.
- 4 Compute bin counts $C_j = \sum_{i=1}^n \mathbb{I}[\hat{w}_i \in B_j]$
- 5 Let $\tilde{S}_1, \dots, \tilde{S}_{(1+2s)/\theta} = \text{PrivatePartialSums}\left([C_j]_{j=1}^{(1+2s)/\theta}; \psi\right)$ (see Appendix C.1)
- 6 Let $j^* = \min \left\{ j \mid \tilde{S}_j + \frac{1}{\sqrt{\psi}} \left(1 + \frac{\log(\frac{1+2s}{\theta}) + \gamma_E}{\pi} \right) \left(\sqrt{\log \frac{1+2s}{\theta}} + \sqrt{\log \frac{2}{\beta}} \right) \geq k \right\}$
- 7 **Return** the triple $(\hat{w}_i, \ell_{j^*}, \mathbb{I}[\hat{w}_i \leq \ell_{j^*}])$ to each individual $i \in [n]$

► **Theorem 12** (Conditions for Algorithm 1 to achieve low regret). *Fix ψ, θ, s , and k , and let ℓ_{j^*} be the threshold returned by Algorithm 1. Define the random variable $Y_{i,j} = \mathbb{I}[\hat{w}_i \in B_j]$. If (1) $\forall j$, $\{Y_{i,j}\}_{i=1}^P$ are jointly independent and (2) $\exists p \in [0, 1]$ such that $\forall i, j$, $\mathbb{E}[Y_{i,j}] \leq p$, then with probability at least $1 - \beta/2$, $k - \sum_{i=1}^P \mathbb{I}[\hat{w}_i \leq \ell_{j^*}]$ is at most:*

$$ks + np + \left(\frac{2}{\sqrt{\psi}} \left(1 + \frac{\log(\frac{1+2s}{\theta}) + \gamma_E}{\pi} \right) + \sqrt{np} \right) \left(\sqrt{\log \frac{1+2s}{\theta}} + \sqrt{\log \frac{2}{\beta}} \right) \quad (6)$$

Combining the preceding two results, we have that with probability at least $1 - \beta$, Algorithm 1 simultaneously stays under budget and satisfies our regret bound. The next two corollaries combine Theorem 12 and Lemma 4 to bound the regret of Algorithm 1 against both stochastic and oblivious adversaries when run on the full population:

► **Corollary 13** (Stochastic adversary regret). *If $s = 0$ and welfare scores are sampled i.i.d. from a distribution with maximum density γ , then the conditions of Theorem 12 are satisfied with $p = \gamma\theta$. Under these assumptions, choosing $\theta = \frac{1}{P\pi\sqrt{\psi}}$ yields:*

$$\overline{\text{REGRET}}(\widetilde{\text{ILA}}, w; k) \leq (1 + o(1)) \cdot \left[\frac{2 \log^{3/2} P + \gamma + 2 \log P \sqrt{\log(2/\beta)}}{\pi \sqrt{\psi}} \right] \quad (7)$$

► **Corollary 14** (Oblivious adversary regret). *If $s > 0$ and w is arbitrary, then the conditions of Theorem 12 are satisfied with $p = \frac{\theta}{2s}$. Choosing $s = \frac{1}{k\pi\sqrt{\psi}}$ and $\theta = \frac{2s \log^{3/2} P}{P\pi\sqrt{\psi}}$ yields:*

$$\overline{\text{REGRET}}(\widetilde{\text{ILA}}, w; k) \leq (1 + o(1)) \cdot \left[\frac{3 \log^{3/2} P + 2 \log P \sqrt{\log(2/\beta)}}{\pi \sqrt{\psi}} \right] \quad (8)$$

3.2 ULA

Our strategy for privatizing ULA is comparatively straightforward. The replace-one sensitivity of the unit profile vector ρ is $1/N$, and so by Lemma 31, adding Gaussian noise with variance $\sigma^2 = (2\psi N^2)^{-1}$ suffices to satisfy ψ -zCDP. Denote the private unit profile vector by $\tilde{\rho}$. Using Lemma 6 alongside standard results about the mean and concentration of various norms of Gaussian vectors, we derive the following unified bound:

Algorithm 2 ψ -zCDP ULA with Public Unit Membership.

Input:

- Vector of welfare scores $w = [w_1, \dots, w_n] \subseteq [0, 1]^n$
- Partition of $[n]$ into units $\mathcal{U} = \{U_1, \dots, U_M\}$; privacy parameter $\psi > 0$; aid budget k

- 1 Sample $\tilde{\rho}_j \sim \mathcal{N}(\rho_j, \frac{1}{2|U_j|\psi})$; Initialize $x_1, \dots, x_M = 0$
- 2 **while** $k > 0$ **do**
- 3 Choose $j = \arg \min \{\tilde{\rho}_j \mid U_j \in \mathcal{U}\}$
- 4 Allocate $x_j = \min(|U_j|, k)$;
- 5 Update budget $k = k - x_j$; Update remaining units $\mathcal{U} = \mathcal{U} \setminus \{U_j\}$
- 6 **end**
- 7 **Return** The unit-level allocation $(x_1, \dots, x_M)$

► **Theorem 15.** Fix $\psi > 0$. Then, with probability at least $1 - \beta$, we have that:

$$\overline{\text{REGRET}}(\widetilde{\text{ULA}}, w; k) \leq k\bar{\rho}_K + \frac{1}{N\sqrt{\psi}} \cdot \min \begin{cases} 2 \min(k, P-k) \sqrt{\log(2M)} + \sqrt{\log(1/\beta)} \\ \sqrt{\min(k, P-k)P} + \sqrt{\log(1/\beta)} \\ \frac{P}{\sqrt{\pi}} + \sqrt{M \log(1/\beta)} \end{cases} \quad (9)$$

4 Accounting for the Cost of Sampling Welfare Scores

We next consider a setting where the administrator must pay an upfront cost of $c_2 n$ to observe the welfare of n individuals, reflecting the cost of e.g. running an RCT or conducting a survey. This is in on top of the baseline cost of ck' to allocate aid to k' individuals, which introduces an exploration/exploitation-style tradeoff: we can improve the efficiency by collecting more data, but because our budget is fixed, this comes at the cost of aiding fewer people.

4.1 ILA

For the remainder of this section, we fix $k = B/c$ and define the ratio $\lambda = c_2/c$. Our goal is to select parameters n, k' to minimize our regret relative to the best k -allocation, subject to the budget constraint that $ck' + c_2 n = B \Rightarrow k' + \lambda n = k$. We derive matching upper and lower bounds for this problem, uncovering a sharp phase-transition: depending on how large λ is, it is optimal to either allocate uniformly at random, or else sample at least $k/2$ people.

► **Theorem 16.** There exists a strategy for ILA with sampling with worst-case expected regret

$$\overline{\text{REGRET}}(\text{ILA}_\lambda, w; k) \leq \begin{cases} \frac{P\lambda k}{P\lambda + k} + \sqrt{\frac{\min(Pk, P^2\lambda)}{16(P\lambda + k)}} & \lambda < \frac{P-k}{P} \\ k \left(1 - \frac{k}{P}\right) & \lambda \geq \frac{P-k}{P} \end{cases} \quad (10)$$

Moreover, any strategy for ILA with sampling must suffer worst-case regret of at least:

$$\overline{\text{REGRET}}(\text{ILA}_\lambda, w; k) \geq \begin{cases} \frac{(P\lambda - \frac{1}{4})k}{P\lambda + k} & \lambda < \frac{P-k}{P} \\ k \left(1 - \frac{k}{P}\right) & \lambda \geq \frac{P-k}{P} \end{cases} \quad (11)$$

Combining Theorem 16 and Corollary 14 yields the following corollary, which summarizes the utility guarantees of private ILA in this setting:

► **Corollary 17.** Let $\widetilde{\text{ILA}}_\lambda$ be the algorithm that samples $n = \frac{Pk}{P\lambda+k}$ scores and then uses Algorithm 1 to approximately allocate aid to the people with the $k' = \frac{k^2}{P\lambda+k}$ lowest scores in the sample. Then M satisfies Joint ψ -zCDP and with probability at least $1 - \beta$, $\overline{\text{REGRET}}(\widetilde{\text{ILA}}_\lambda, w; k)$ is at most:

$$\frac{P \left(\lambda k + \sqrt{\lambda k \log(3/\beta)} \right)}{P\lambda + k} + \sqrt{\frac{\min(Pk, P^2\lambda)}{16(P\lambda + k)}} + (1 + o(1)) \left[\frac{3 \log^{3/2} P + 2 \log P \sqrt{\log(3/\beta)}}{\pi \sqrt{\psi}} \right]$$

4.2 ULA

We begin by conducting a stratified sample of n individuals across units, giving us n/M observations per unit which can be provided as input to Algorithm 2. Our final error is a combination of the statistical error from sampling and the error from additive noise.

An interesting feature of this problem (also observed in [16]) is that the worst case in a statistical sense would occur if $\rho_j = 0.5$ for all j , but in this case, the excess regret of our *ranking* would be 0. Building on this intuition, we show that for any fixed vector ρ , the average variance across all components must be small whenever inequality is high:

► **Lemma 18.** Let G be the Gini coefficient of the unit profile vector ρ . Then we have that $\frac{1}{M} \sum_{j=1}^M \rho_j(1 - \rho_j) \leq \bar{\rho}_M(1 - \bar{\rho}_M) - 3\bar{\rho}_M^2 \cdot G^2$

Combining this result with Lemma 6 yields the following theorem:

► **Theorem 19.** Fix k, λ , and ρ . Then $\overline{\text{REGRET}}(\widetilde{\text{ULA}}_\lambda, w; k)$ is at most:

$$\bar{\rho}_K k + (1 - \bar{\rho}_K) \lambda n + \min \left\{ \begin{aligned} &\sqrt{\frac{M^2}{2\psi n^2} + \frac{P-n}{4nN}} \cdot 2 \min(k, P-k) \sqrt{2 \log(2M)} \\ &\sqrt{\frac{M^2}{2\psi n^2} + \frac{(P-n)(\bar{\rho}_M(1-\bar{\rho}_M) - 3\bar{\rho}_M^2 G^2)}{nN}} \cdot \sqrt{2 \min(k, P-k) P} \\ &\sqrt{\frac{M^2}{2\psi n^2} + \frac{(P-n)(\bar{\rho}_M(1-\bar{\rho}_M) - 3\bar{\rho}_M^2 G^2)}{nN}} \cdot P \sqrt{\frac{2}{\pi}} \end{aligned} \right. \quad (12)$$

► **Corollary 20.** Let $C = \min \left(2 \min(k, P-k) \sqrt{2 \log(2M)}, \sqrt{2 \min(k, P-k) P}, P \sqrt{\frac{2}{\pi}} \right)$. Choosing $n = \max \left(\left(\frac{CM}{\lambda(2\psi)^{1/2}} \right)^{1/2}, \left(\frac{C^2 M}{16\lambda^2} \right)^{1/3} \right)$, the marginal increase in the expected normalized regret of $\widetilde{\text{ULA}}_\lambda$ over non-private ULA is at most: $2^{3/4} (CM\lambda/\sqrt{\psi})^{1/2} + 2^{-4/3} (C^2 M\lambda)^{1/3}$

4.3 Asymptotic Regimes

For both ULA and ILA, the price of privacy is asymptotically negligible provided that $k = \Theta(P)$. Understanding their asymptotic performance therefore boils down to a comparison of the linear terms in each strategy's regret bound. By combining our previous results, we can define three possible asymptotic regimes in terms of the key parameters G, λ , and $\bar{\rho}_M$:

- **ULA > ILA > RAND.** This occurs whenever $\frac{k}{P} \cdot \frac{\bar{\rho}_M(1-G)}{1-\bar{\rho}_M(1-G)} < \lambda < 1 - \frac{k}{P}$; sampling is cheap enough for ILA to be more effective than random allocation, but too expensive to justify the marginal cost of needing to sample a larger percentage of the population.
- **ULA > ILA = RAND.** This occurs whenever $1 - \frac{k}{P} \leq \lambda$ and $G > 0$; sampling is too expensive to meaningfully improve the efficiency of ILA, but ULA is still able to take advantage of the information conveyed by unit membership.
- **ILA $\geq$ ULA $\geq$ RAND.** This occurs whenever $\lambda < \min(1 - \frac{k}{P}, \frac{k}{P} \cdot \frac{\bar{\rho}_M(1-G)}{1-\bar{\rho}_M(1-G)})$; sampling is sufficiently cheap for ILA to make efficient use of it, while low inequality and high average welfare make unit membership a weak signal for identifying the worst-off.

A particularly interesting consequence of the above is that ULA always dominates ILA when $\lambda \geq 1$, which can be interpreted as representing settings where the only efficient means to measure welfare (or treatment effects) is to actually allocate aid and observe the result.

5 Learning Welfare Scores from Auxiliary Information

In the final setting we consider, we can once again pay a cost of $c_2 = \lambda c$ to observe a single individual's welfare. But unlike in the previous setting, every individual also has a *feature vector* $x_i \in \mathcal{X}$ and a binary label $y_i = \mathbb{I}[w_i > 1 - \delta_w]$ drawn i.i.d. from a joint distribution $\mathcal{D}$ over $\mathcal{X} \times \{0, 1\}$. We assume that features (but not labels) are already known to the administrator, whose goal is to learn a convex model to predict $\eta := \mathbb{E}[y | x]$ subject to DP, with both labels and features considered private. For this purpose, we make use of Algorithm 2 of Feldman et al. [21], which achieves optimal rates on population risk (see Theorem 32).

Our final allocation is computed using only the data of individuals who weren't included in the original sample. This allows us to use parallel composition [48], and because both ULA and ILA sample $o(P)$ welfare scores, it induces only asymptotically negligible regret.

For the purposes of relating our model's risk to the downstream efficiency of the induced allocation, it is useful to additionally require that the loss functions we use are *proper* [14, 8, 50], meaning that $\mathcal{L}(\theta)$ is minimized when $\hat{\eta} := f_\theta(x) = \eta$. In the sequel, we will focus primarily on the squared loss $f(\theta, x, y) = (y - f_\theta(x))^2$, a canonical proper loss function.

We define regret in this stochastic setting against the best allocation *in hindsight*, i.e. $\mathbb{E}_{w \sim \mathcal{D}^P}[\max_{|\mathcal{J}|=k} \text{VAL}(\mathcal{J}, w)]$. As a consequence, ILA with a Bayes' optimal model could still suffer significant regret if the variance of $y|x$ is high. It will emerge that ULA is much more robust to this sort of unpredictability, which we formalize with the following definition:

► **Definition 21** (Excess-Risk Decomposition). *Let ℓ denote the squared loss, and let $\mathcal{L}(\theta) := \mathbb{E}[\ell(\theta, x, y)] = \alpha$. Then, if $x, y \sim \mathcal{D}$,*

$$\alpha = \mathbb{E}[(f_\theta(x) - y)^2] = \underbrace{\mathbb{E}[(f_\theta(x) - \mathbb{E}[y | x])^2]}_E + \underbrace{\mathbb{E}[(y - \mathbb{E}[y | x])^2]}_{\sigma^2} \quad (13)$$

Where E is the excess risk, and σ^2 is the irreducible or Bayes' risk. The excess risk of the optimal classifier $\theta^* = \arg \min_{\theta \in \Theta} \mathcal{L}(\theta)$ is denoted by E^* .

5.1 ILA

ILA proceeds by first learning a model to predict the probability that an individual is low-welfare, and then allocating aid according to the predicted probabilities in ascending order. To start out, we take the expected loss of the learned model as a constant, and then later incorporate the error bounds from Theorem 32 to derive the optimal choice of n .

► **Lemma 22.** *Let f be a binary classifier with expected ℓ_2^2 loss α , and let $\widehat{ILA}$ denote the algorithm that allocates aid to individuals in ascending order based on $\hat{y}_i = f(x_i)$. Then with probability at least $1 - \beta$, we have:*

$$\overline{\text{REGRET}}(\widehat{ILA}, w; k) \leq k' \cdot \bar{\eta}_{k'} + P\sqrt{\alpha} + P^{3/4}\sqrt{\log(1/\beta)/2} \quad (14)$$

► **Theorem 23.** Let $\widetilde{ILA}$ be the algorithm that samples n individuals from the population uniformly at random, fits a model $\hat{\theta}$ with expected ℓ_2^2 error α subject to ψ -zCDP, and then runs Algorithm 1 on the estimated scores of the remaining population with privacy parameter ψ . Then $\widetilde{ILA}$ satisfies Joint ψ -zCDP, and:

$$\overline{\text{REGRET}}(\widetilde{ILA}, w; k) \leq n\lambda + k'\bar{\eta}_{k'} + P\sqrt{\alpha} + O\left(\frac{\log^{3/2} P}{\sqrt{\psi}}\right) \quad (15)$$

► **Corollary 24.** When $\widetilde{ILA}$ is instantiated with Algorithm 2 of [21], its asymptotic normalized regret grows like:

$$n\lambda + k'\bar{\eta}_{k'} + P\sqrt{\sigma^2 + E^* + 10LD\left(\frac{1}{\sqrt{n}} + \frac{\sqrt{p}}{n\sqrt{2\psi}}\right)} + O\left(\frac{\log^{3/2} P}{\sqrt{\psi}}\right) \quad (16)$$

Choosing $n = \max\left(\left(\frac{P}{\lambda}\sqrt{10LD(2\psi)^{-1/2}\sqrt{p}}\right)^{2/3}, \left(\frac{P}{\lambda}\sqrt{10LD}\right)^{4/5}\right)$, we can guarantee a worst-case regret bound of:

$$k'\bar{\eta}_{k'} + P\sqrt{\mathcal{L}(\theta^*)} + 2\lambda^{1/3}\left(\frac{P\sqrt{10LD\sqrt{p}}}{(2\psi)^{1/4}}\right)^{2/3} + 2\lambda^{1/5}\left(P\sqrt{10LD}\right)^{4/5} \quad (17)$$

If we additionally have access to a lower bound on the risk of the optimal classifier $\mathcal{L}(\theta^*) \geq L^\downarrow$, then we can instead choose $n = \max\left(\sqrt{\frac{5PLD\sqrt{p}}{\lambda\sqrt{2\psi}L^\downarrow}}, \left(\frac{5PLD}{\lambda\sqrt{L^\downarrow}}\right)^{2/3}\right)$ for an instance-specific regret bound of:

$$k'\bar{\eta}_{k'} + P\sqrt{\mathcal{L}(\theta^*)} + 2\sqrt{\frac{5PLD\lambda\sqrt{p}}{\sqrt{2\psi}L^\downarrow}} + 2\left(\frac{5PLD\lambda}{\sqrt{L^\downarrow}}\right)^{2/3} \quad (18)$$

5.2 ULA

In the previous settings we've considered, unit membership could be thought of as a particularly simple feature vector corresponding the standard basis of $\mathbb{R}^M$. We will generalize this idea by assuming we are given a partition $\mathcal{S}$ over the feature space $\mathcal{X} \subseteq \mathbb{R}^p$, which induces a partition over individuals into $|\mathcal{S}|$ units. This modeling decision is motivated by the sorts of bureaucratically legible and overlapping categories that are frequently used to target aid in existing systems [33], such as “disabled veterans.” We treat this choice of categories as exogenous and do not consider the problem of trying to *learn* an efficient partition.

One minor complication is that up until this point, we have assumed that the unit membership of each individual is publicly known, which may be inappropriate if units correspond to sensitive traits like poverty or disability rather than e.g. geographic regions. We therefore introduce Algorithm 3, a Joint DP variant Algorithm 2 that can accommodate private unit membership at essentially no cost, provided that no unit is too small.

► **Lemma 25.** Choosing $\frac{\psi_2}{\psi_1} = \left(\frac{3\log^{3/2} P}{\pi} \cdot \frac{N}{C\sqrt{2}}\right)^{2/3}$, the expected normalized regret of Algorithm 3 is at most:

$$k\bar{\rho}_K + \frac{1}{\sqrt{\psi}}\left[\left(\frac{3}{\pi}\right)^{2/3}\log P + \left(\frac{C\sqrt{2}}{N}\right)^{2/3}\right]^{3/2} = k\bar{\rho}_K + O\left(\frac{C}{N\sqrt{\psi}}\right) \quad (19)$$

With C defined as in Corollary 20, which matches the rate achieved by Algorithm 2 in the public unit membership case up to constant factors.

Algorithm 3 Joint ψ -zCDP ULA with Private Unit Membership.

Input:

- Vector of welfare scores $w = [w_1, \dots, w_P] \subseteq [0, 1]^P$
 - Partition of $[P]$ into units $\mathcal{U} = \{U_1, \dots, U_M\}$ with minimum unit size N
 - Privacy parameters $\psi_1 + \psi_2 = \psi$; aid budget k
- 1 Sample $\tilde{\rho}_j \sim \mathcal{N}(\rho_j, \frac{1}{N^2\psi_1})$;
 - 2 For each individual $i \in [P]$ with corresponding unit j , set $\tilde{w}_i = \tilde{\rho}_j$
 - 3 **Return** output of Algorithm 1 with input $\tilde{w}, \psi_2$ and s, θ set as in Corollary 14
-

Because the irreducible error term from Definition 21 is independent and zero-mean, ULA is able to effectively ignore it when taking averages across many individuals in a single unit, which is reflected in the following lemma:

► **Lemma 26.** *Let f be a binary classifier with expected ℓ_2^2 loss α , and let $\widehat{ULA}$ denote the algorithm that allocates aid to units in ascending order based on $\hat{\rho}_j = \frac{1}{N_j} \sum_{i \in U_j} f(\hat{x}_i)$. Then in expectation, we have:*

$$\overline{\text{REGRET}}(\widehat{ULA}, w; k) \leq k\bar{\rho}_K + \sqrt{P|\mathcal{P}|\sigma^2} + P\sqrt{\alpha - \sigma^2} \quad (20)$$

► **Theorem 27.** *Let $\widetilde{ULA}$ be the algorithm that samples n individuals from the population uniformly at random, fits a model $\hat{\theta}$ with expected ℓ_2^2 error α subject to ψ -zCDP, and then runs Algorithm 3 on the estimated welfare scores of the remaining population with privacy parameter ψ and partition $\mathcal{P}$. Then $\widetilde{ULA}$ satisfies Joint ψ -zCDP in the replace-one adjacency definition with private unit membership, and:*

$$\overline{\text{REGRET}}(\widetilde{ULA}, w; k) \leq k\bar{\rho}_K + P\sqrt{\alpha - \sigma^2} + \sqrt{P|\mathcal{P}|\sigma^2} + O\left(\frac{C}{N\sqrt{\psi}}\right) \quad (21)$$

Incorporating the cost of sampling and the utility guarantees from Theorem 32 yields the following corollary:

► **Corollary 28.** *When $\widetilde{ULA}$ is instantiated with Algorithm 2 of [21], its asymptotic normalized regret grows like:*

$$k\bar{\rho}_K + (1 - \bar{\rho}_K)n\lambda + \sqrt{P|\mathcal{P}|\sigma^2} + P\sqrt{E^* + 10LD\left(\frac{1}{\sqrt{n}} + \frac{\sqrt{p}}{n\sqrt{2\psi}}\right)} + O\left(\frac{C}{N\sqrt{\psi}}\right) \quad (22)$$

Choosing $n = \max\left(\left(\frac{P}{\lambda}\sqrt{10LD(2\psi)^{-1/2}\sqrt{p}}\right)^{2/3}, \left(\frac{P}{\lambda}\sqrt{10LD}\right)^{4/5}\right)$, we can guarantee a worst-case regret bound of:

$$k\bar{\rho}_K + P\sqrt{E^*} + \sqrt{P|\mathcal{P}|\sigma^2} + O\left(\frac{C}{N\sqrt{\psi}}\right) \quad (23)$$

$$+ (2 - \bar{\rho}_K) \left[\left(\frac{P\sqrt{10LD\lambda\sqrt{p}}}{(2\psi)^{1/4}} \right)^{2/3} + \lambda^{1/5} \left(P\sqrt{10LD} \right)^{4/5} \right] \quad (24)$$

If we additionally have access to a lower bound on the excess risk of the optimal classifier $E^* \geq E^\downarrow > 0$, then we can instead choose $n = \max\left(\sqrt{\frac{5PLD\sqrt{p}}{\lambda\sqrt{2\psi}E^\downarrow}}, \left(\frac{5PLD}{\lambda\sqrt{E^\downarrow}}\right)^{2/3}\right)$ for an instance-specific regret bound of:

$$k\bar{\rho}_K + P\sqrt{E^*} + \sqrt{P|\mathcal{P}|\sigma^2} + (2 - \bar{\rho}_K) \left[\sqrt{\frac{5PLD\lambda\sqrt{p}}{\sqrt{2\psi}E^\downarrow}} + \left(\frac{5PLD\lambda}{\sqrt{E^\downarrow}} \right)^{2/3} \right] + O\left(\frac{C}{N\sqrt{\psi}}\right) \quad (25)$$

5.3 Asymptotic Regimes

As in Section 4, the price of privacy is asymptotically negligible for both ILA and ULA, and so it suffices to compare the leading linear terms. For ULA, our regret grows like $k\bar{\rho}_K + P\sqrt{E^*} \leq k(1-G)\bar{\rho}_M + P\sqrt{E^*}$, while for ILA it grows like $P\sqrt{E^* + \sigma^2} + k'\bar{\eta}_{k'} \geq P\sqrt{E^* + \sigma^2}$. Using the first-order approximation $\sqrt{E^*} = \sqrt{E^* + \sigma^2 - \sigma^2} \leq \sqrt{E^* + \sigma^2} - \frac{\sigma^2}{2\sqrt{E^* + \sigma^2}}$, we obtain the following sufficient condition for ULA to outperform ILA asymptotically:

$$k(1-G)\bar{\rho}_M \leq \frac{P\sigma^2}{2\sqrt{E^* + \sigma^2}} \quad (26)$$

$$2\frac{k}{P}(1-G)\bar{\rho}_M\sqrt{\mathcal{L}(\theta^*)} \leq \sigma^2 \quad (27)$$

If the special case where Θ contains the Bayes' optimal classifier, $\mathcal{L}(\theta^*) = \sigma^2$ and the condition simplifies to:

$$\left(\frac{1}{\sigma}\right)\left(\frac{k}{P}\right)(1-G)\bar{\rho}_M \leq \frac{1}{2} \quad (28)$$

Each term on the left-hand side has a nice interpretation: $1/\sigma$ represents the intrinsic predictability of the outcomes we care about, k/P represents the size of our budget, $1-G$ represents how *equally* welfare is distributed across units, and $\bar{\rho}_M$ represents the average welfare across units. Taken as a whole, we see that ULA dominates ILA asymptotically whenever the product of these four terms is not too large.

5.4 The Value of Modeling

Since training a model to predict welfare implicitly requires the ability to sample welfare scores from the population, it would be possible to dispense with the modeling step entirely and use the algorithms from Section 4 even in this stochastic setting. We therefore consider one final comparison between the regret bounds obtained for each strategy in this section and those for the same strategy in Section 4.

In the case of ILA, our linear regret term in the pure sampling setting scaled like $\frac{P\lambda k}{P\lambda + k}$, while the regret of the model-based approach scaled like $P\sqrt{\mathcal{L}(\theta^*)}$. Modeling therefore emerges as a stronger strategy in settings where λ is relatively large. In particular, we saw that sampling-based strategies could do no better than random allocation for $\lambda > 1 - \frac{k}{P}$, while modeling-based strategies can in principle achieve near-optimal regret even for $\lambda > 1$.

In the case of ULA, we ignore the common linear term of $k\bar{\rho}_K$. Then, the excess regret of the sampling based approach grows like $O(P^2 M \lambda)^{1/3}$, while the excess regret of the modeling approach grows like $\lambda^{1/5} P^{4/5} + P\sqrt{E^*}$. We can therefore identify three natural regimes of interest:

- If $E^* > 0$ and $M = o(P)$, then the modeling approach can *never* outperform pure sampling asymptotically. The key reason for this is that the error from sampling is purely driven by variance, whereas the error from modeling can involve bias as well.
- If $E^* = 0$, then modeling can outperform sampling asymptotically only when $M = \Omega((P/\lambda)^{2/5})$. This is because when M is very large, the modeling based approach can leverage data from similar individuals across units to avoid sampling every unit directly.
- If $E^* > 0$ and $M = \Theta(P)$, then it is possible for modeling to outperform sampling if $\lambda \gtrsim (E^*)^{3/2}$, which once again occurs because modeling allows us to avoid sampling many individuals from every unit at the cost of introducing some bias.

It is interesting to note that the conditions for modeling to improve the performance of ULA seem to be considerably stricter than for ILA. One possible interpretation for this fact is that unit-level targeting in itself already constitutes a simple sort of model, and so the ability to leverage information about one individual to infer the welfare of another was already available to ULA even before making any distributional assumptions. In other words, ULA has less to gain than ILA from additional tools for making such inferences.

6 Discussion

We conclude our analysis by reflecting on the strengths and limitations of our modeling choices and suggesting possible avenues for future work.

6.1 On Our Modeling Choices

6.1.1 Our Choice of Social Welfare Ordering

For both ULA and ILA, the objective function we seek to maximize is $\sum_{i=1}^P w_i$, which corresponds to the classical Utilitarian social welfare ordering [49]. One limitation of this ordering is that it is indifferent to *distributive* concerns; a world where half of the population has welfare 1 and the other half has welfare 0 is deemed just as desirable as a world where everybody has welfare 0.5. Possible alternatives that are more sensitive to distributive questions include the Nash social welfare ordering, which corresponds to maximizing $\prod_{i=1}^P w_i$, and the egalitarian social welfare ordering, which (roughly) seeks to maximize $\min_i w_i$.

The main challenge with both of these alternative objectives in the context of DP is that they can be extremely sensitive to changes in a single individual's welfare; one person with welfare 0 forces both objective functions to be 0 regardless of the remaining data. In contrast, the Utilitarian objective is a sum of bounded data points, which is a well-studied object in the DP literature. We therefore believe that the Utilitarian objective is the natural starting point for our investigation, and leave the design of DP algorithms for optimizing either of the other objectives (or choosing between multiple equal-value allocations [32]) to future work.

6.1.2 Statistical Estimation vs. Pure Selection

The strategies that we consider all derive their allocations from privatized estimates of important statistics, such as the empirical CDF of welfare scores or the unit profile vector. Outputting a private aid allocation is fundamentally a *selection* task, however, and so one could argue that it would be more efficient to skip the estimation step and use methods, like those contained in [24, 7, 60, 34], which optimize over the set of possible allocations directly. In the specific domain of aid allocation, however, we believe that approaches based on statistical estimation have major advantages in transparency: in addition to outputting an allocation, our algorithms can release a privatized statistic which justifies why that allocation is fair. A further benefit is that the intermediate statistics we compute are general enough to be repurposed for other tasks without expending any additional privacy budget.

6.1.3 What Our Model Does Not Capture

Our model assumes a very simple and schematic relationship between welfare and treatment effects which may not hold exactly in real settings. The authors of [58] on whose work we build upon consider more complicated treatment effects and show that when $\tau(w)$ is Lipschitz and concave, and when the distribution of welfare scores is sufficiently smooth, estimating welfare

alone is still sufficient to derive efficient allocations. An alternative approach, which others have advocated for on philosophical grounds, is to reframe the problem in terms of predicting treatment effects directly [9, 43]. To this end, there is a large body of works, spanning many sciences, which study how heterogeneous treatment effects can be learned [61, 42], including under DP constraints [51]. Although we do not directly pursue this point, we remark that one virtue of the simplicity of our model is that our results can be easily restated to take either treatment effects or welfares as primary.

6.2 Conclusion and Future Work

In this work, we separately considered two possible constraints on data availability and found that in all cases, both ULA and ILA could be modified to satisfy DP without affecting their asymptotic regret. In Section 4, we imagined that welfare scores were fixed arbitrarily and designed robust algorithms for achieving low regret in the worst-case. Here, the key difference between ILA and ULA was that ULA required only a sublinear number of samples to achieve nearly-optimal targeting, while ILA was obliged to expend a constant fraction of its budget on sampling. As a result, the optimal choice of strategy depended strongly on λ , the relative cost of measuring welfare compared to the cost of treatment.

Meanwhile, in Section 5, we assumed that welfare scores and features were sampled i.i.d. from some fixed joint distribution and designed learning-based algorithms that could estimate unseen welfare scores with high accuracy. Unlike the previous setting, we found that there was essentially no difference between the sampling required by ILA vs. ULA; instead, the choice of optimal strategy depended on a new quantity, σ^2 , representing the fundamental unpredictability of welfare scores conditioned on the available data. Because ULA was able to “wash out” this uncertainty by averaging many measurements together, it was significantly more robust than ILA in settings where outcomes are fundamentally difficult to predict.

A natural next step (with or without DP) is to combine these two settings as follows: imagine an administrator trying to select an allocation strategy. To start out, they have access to relatively coarse features that encode unit membership and nothing else. But by paying some cost c_p , they can obtain p new features for a single individual – for instance, performance on standardized assessments [31] or proxies for socioeconomic status [41]. By expanding their model class in this way, the administrator can potentially decrease the irreducible error of their prediction problem, but this will come at the cost of increasing their estimation error and obliging them to spend more of their budget on acquiring extended feature vectors. Seen in this light, ILA and ULA might appear less as diametric opposites, and more as two points along a continuum of tradeoffs between the error arising from coarse targeting, sampling, and epistemic uncertainty.

An alternative path to extend our model is to incorporate adaptive or strategic behavior. On the administrative side, this might involve conducting repeated rounds of sampling and/or intervention, which could enable meaningful reductions in sample complexity using techniques from the bandits literature [6, 34] or permit treatment effects and allocations to be learned jointly [36, 64]. On the population side, this might involve considering the possibility of strategic behavior from both individuals [13] and organized groups [26] wishing to influence the administrator’s decision. It would be particularly interesting to extend our model to incorporate insights from the large body of literature on *performative prediction* [59, 46, 27, 25], whereby forecasts can influence the very events they predict, and to investigate how the stability induced by DP impacts these dynamics.

References

- 1 Abdul Latif Jameel Poverty Action Lab (J-PAL). The challenges of targeting social protection programs. *African Arguments*, March 2022. Blog post. URL: <https://www.povertyactionlab.org/blog/3-10-22/challenges-targeting-social-protection-programs>.
- 2 Robert J Adler and Jonathan E Taylor. Gaussian inequalities. *Random Fields and Geometry*, pages 49–64, 2007.
- 3 Emily Aiken, Suzanne Bellue, Dean Karlan, Chris Udry, and Joshua E Blumenstock. Machine learning and phone data can improve targeting of humanitarian aid. *Nature*, 603(7903):864–870, 2022. doi:10.1038/s41586-022-04484-9.
- 4 Emily L. Aiken. *Targeting Social Protection Programs with Machine Learning and Digital Data*. PhD thesis, University of California Berkeley, USA, 2024. URL: <https://www.escholarship.org/uc/item/0cj4c266>.
- 5 Emily L. Aiken, Guadalupe Bedoya, Joshua Blumenstock, and Aidan Coville. Program targeting with machine learning and mobile phone data: Evidence from an anti-poverty intervention in afghanistan. *CoRR*, abs/2206.11400:103016, 2022. doi:10.48550/arXiv.2206.11400.
- 6 Jean-Yves Audibert, Sébastien Bubeck, and Rémi Munos. Best arm identification in multi-armed bandits. In Adam Tauman Kalai and Mehryar Mohri, editors, *COLT 2010 - The 23rd Conference on Learning Theory, Haifa, Israel, June 27-29, 2010*, pages 41–53. Omnipress, 2010. URL: <http://colt2010.haifa.il.ibm.com/papers/COLT2010proceedings.pdf#page=49>.
- 7 Mitali Bafna and Jonathan Ullman. The price of selection in differential privacy. In *Conference on Learning Theory*, pages 151–168. PMLR, 2017. URL: <http://proceedings.mlr.press/v65/bafna17a.html>.
- 8 Han Bao. Proper losses, moduli of convexity, and surrogate regret bounds. In *The Thirty Sixth Annual Conference on Learning Theory*, pages 525–547. PMLR, 2023. URL: <https://proceedings.mlr.press/v195/bao23a.html>.
- 9 Chelsea Barabas, Madars Virza, Karthik Dinakar, Joichi Ito, and Jonathan Zittrain. Interventions over predictions: Reframing the ethical debate for actuarial risk assessment. In *Conference on fairness, accountability and transparency*, pages 62–76. PMLR, 2018. URL: <http://proceedings.mlr.press/v81/barabas18a.html>.
- 10 Daniel Björkegren and Darrell Grissen. Behavior revealed in mobile phone usage predicts loan repayment. *CoRR*, abs/1712.05840(3):618–634, 2017. doi:10.48550/arXiv.1712.05840.
- 11 Joshua E. Blumenstock and Nitin Kohli. Big data privacy in emerging market fintech and financial services: A research agenda. *CoRR*, abs/2310.04970, 2023. doi:10.48550/arXiv.2310.04970.
- 12 Mary Bruce, John M Bridgeland, Joanna Hornig Fox, and Robert Balfanz. On track for success: The use of early warning indicator and intervention systems to build a grad nation. *Civic Enterprises*, 2011.
- 13 Michael Brückner and Tobias Scheffer. Stackelberg games for adversarial prediction problems. In *Proceedings of the 17th ACM SIGKDD international conference on Knowledge discovery and data mining*, pages 547–555, 2011. doi:10.1145/2020408.2020495.
- 14 Andreas Buja, Werner Stuetzle, and Yi Shen. Loss functions for binary class probability estimation and classification: Structure and applications. *Working draft, November*, 3:13, 2005.
- 15 Mark Bun and Thomas Steinke. Concentrated differential privacy: Simplifications, extensions, and lower bounds. In *Theory of cryptography conference*, pages 635–658. Springer, 2016. doi:10.1007/978-3-662-53641-4_24.
- 16 Sílvia Casacuberta and Moritz Hardt. Good allocations from bad estimates. *arXiv preprint*, 2026. arXiv:2601.05597.
- 17 Du Chen and Geoffrey A. Chua. Noisy dual mirror descent: A near optimal algorithm for jointly-dp convex resource allocation. In Amir Globersons, Lester Mackey, Danielle Belgrave, Angela Fan, Ulrich Paquet, Jakub M. Tomczak, and Cheng Zhang, editors,

- Advances in Neural Information Processing Systems 38: Annual Conference on Neural Information Processing Systems 2024, NeurIPS 2024, Vancouver, BC, Canada, December 10 - 15, 2024*, 2024. URL: http://papers.nips.cc/paper_files/paper/2024/hash/a6e1f6963f65bcc4854691a15460dbd8-Abstract-Conference.html.
- 18 B. R. Data. How Biometric Devices Are Putting Afghans in Danger. URL: <https://interaktiv.br.de/biometrie-afghanistan/en/>.
 - 19 Michael Dinitz, George Z. Li, Quanquan C. Liu, and Felix Zhou. Differentially private matchings. *CoRR*, abs/2501.00926, 2025. doi:10.48550/arXiv.2501.00926.
 - 20 Cynthia Dwork, Frank McSherry, Kobbi Nissim, and Adam Smith. Calibrating noise to sensitivity in private data analysis. In *Theory of cryptography conference*, pages 265–284. Springer, 2006. doi:10.1007/11681878_14.
 - 21 Vitaly Feldman, Tomer Koren, and Kunal Talwar. Private stochastic convex optimization: optimal rates in linear time. In *Proceedings of the 52nd Annual ACM SIGACT Symposium on Theory of Computing*, pages 439–449, 2020. doi:10.1145/3357713.3384335.
 - 22 Unai Fischer-Abaigar, Emily Aiken, Christoph Kern, and Juan Carlos Perdomo. Empirically understanding the value of prediction in allocation, 2026. doi:10.48550/arXiv.2602.08786.
 - 23 Unai Fischer-Abaigar, Christoph Kern, and Juan Carlos Perdomo. The Value of Prediction in Identifying the Worst-Off, January 2025. arXiv:2501.19334 [cs]. doi:10.48550/arXiv.2501.19334.
 - 24 Jennifer Gillenwater, Matthew Joseph, Andres Muñoz Medina, and Mónica Ribero Diaz. A joint exponential mechanism for differentially private top-k. In Kamalika Chaudhuri, Stefanie Jegelka, Le Song, Csaba Szepesvári, Gang Niu, and Sivan Sabato, editors, *International Conference on Machine Learning, ICML 2022, 17-23 July 2022, Baltimore, Maryland, USA*, volume 162 of *Proceedings of Machine Learning Research*, pages 7570–7582. PMLR, PMLR, 2022. URL: <https://proceedings.mlr.press/v162/gillenwater22a.html>.
 - 25 Moritz Hardt, Meena Jagadeesan, and Celestine Mendler-Dünnér. Performative Power, November 2022. arXiv:2203.17232 [cs]. doi:10.48550/arXiv.2203.17232.
 - 26 Moritz Hardt, Eric Mazumdar, Celestine Mendler-Dünnér, and Tijana Zrnic. Algorithmic collective action in machine learning. *CoRR*, abs/2302.04262, 2023. doi:10.48550/arXiv.2302.04262.
 - 27 Moritz Hardt and Celestine Mendler-Dünnér. Performative Prediction: Past and Future, October 2023. arXiv:2310.16608 [cs]. doi:10.48550/arXiv.2310.16608.
 - 28 Monika Henzinger, Jalaj Upadhyay, and Sarvagya Upadhyay. Almost Tight Error Bounds on Differentially Private Continual Counting. *Proceedings of the 2023 Annual ACM-SIAM Symposium on Discrete Algorithms (SODA)*, 2023. doi:10.1137/1.9781611977554.ch183.
 - 29 Justin Hsu, Zhiyi Huang, Aaron Roth, Tim Roughgarden, and Zhiwei Steven Wu. Private Matchings and Allocations. *SIAM Journal on Computing*, 45(6):1953–1984, January 2016. doi:10.1137/15100271X.
 - 30 Justin Hsu, Zhiyi Huang, Aaron Roth, and Zhiwei Steven Wu. Jointly private convex programming. In *Proceedings of the twenty-seventh annual ACM-SIAM symposium on Discrete algorithms*, pages 580–599. SIAM, 2016. doi:10.1137/1.9781611974331.ch43.
 - 31 Gagaoua Ikram, Armelle Brun, and Anne Boyer. A frugal model for accurate early student failure prediction. *CoRR*, abs/2502.00017, 2025. doi:10.48550/arXiv.2502.00017.
 - 32 Shomik Jain, Margaret Wang, Kathleen Creel, and Ashia Wilson. Allocation Multiplicity: Evaluating the Promises of the Rashomon Set, March 2025. arXiv:2503.16621 [cs]. doi:10.48550/arXiv.2503.16621.
 - 33 Rebecca Ann Johnson and Simone Zhang. What is the Bureaucratic Counterfactual? Categorical versus Algorithmic Prioritization in U.S. Social Policy. In *2022 ACM Conference on Fairness, Accountability, and Transparency*, pages 1671–1682, Seoul Republic of Korea, June 2022. ACM. doi:10.1145/3531146.3533223.
 - 34 Marc Jourdan and Achraf Azize. Optimal best arm identification under differential privacy. *CoRR*, abs/2510.17348, 2025. doi:10.48550/arXiv.2510.17348.

- 35 Zoe Kahn, Meyebinesso Farida Carelle Pere, Emily Aiken, Nitin Kohli, and Joshua E Blumenstock. Expanding perspectives on data privacy: Insights from rural togo. *Proceedings of the ACM on Human-Computer Interaction*, 9(2):1–29, 2025. doi:10.1145/3710968.
- 36 Maximilian Kasy and Anja Sautmann. Adaptive treatment assignment in experiments for policy choice. *Econometrica*, 89(1):113–132, 2021.
- 37 Michael Kearns, Mallesh M. Pai, Ryan Rogers, Aaron Roth, and Jonathan Ullman. Robust mediators in large games, 2015. doi:10.48550/arXiv.1512.02698.
- 38 Michael P. Kim and Moritz Hardt. Is Your Model Predicting the Past? In *Equity and Access in Algorithms, Mechanisms, and Optimization*, pages 1–8, Boston MA USA, October 2023. ACM. doi:10.1145/3617694.3623225.
- 39 Jon Kleinberg, Himabindu Lakkaraju, Jure Leskovec, Jens Ludwig, and Sendhil Mullainathan. Human decisions and machine predictions. *The Quarterly Journal of Economics*, 133(1):237–293, 2018.
- 40 Nitin Kohli, Emily L. Aiken, and Joshua Blumenstock. Privacy guarantees for personal mobility data in humanitarian response. *CoRR*, abs/2306.09471(1):28565, 2023. doi:10.48550/arXiv.2306.09471.
- 41 Nitin Kohli and Joshua Blumenstock. Enabling Humanitarian Applications with Targeted Differential Privacy, August 2024. arXiv:2408.13424 [cs]. doi:10.48550/arXiv.2408.13424.
- 42 Sören R Künzle, Jasjeet S Sekhon, Peter J Bickel, and Bin Yu. Metalearners for estimating heterogeneous treatment effects using machine learning. *Proceedings of the national academy of sciences*, 116(10):4156–4165, 2019.
- 43 Lydia T Liu, Serena Wang, Tolani Britton, and Rediet Abebe. Reimagining the machine learning life cycle to improve educational outcomes of students. *Proceedings of the National Academy of Sciences*, 120(9):e2204781120, 2023.
- 44 Ian Lundberg, Rachel Brown-Weinstock, Susan Clampet-Lundquist, Sarah Pachman, Timothy J. Nelson, Vicki Yang, Kathryn Edin, and Matthew J. Salganik. The origins of unpredictability in life trajectory prediction tasks, October 2023. arXiv:2310.12871 [stat]. doi:10.48550/arXiv.2310.12871.
- 45 Eva Luvison, Sylvain Chatel, Justinas Sukaitis, Vincent Graf Narbel, Carmela Troncoso, and Wouter Lueks. A Low-Cost Privacy-Preserving Digital Wallet for Humanitarian Aid Distribution, October 2024. arXiv:2410.15942 [cs]. doi:10.48550/arXiv.2410.15942.
- 46 David Manheim and Scott Garrabrant. Categorizing Variants of Goodhart’s Law, February 2019. arXiv:1803.04585 [cs]. doi:10.48550/arXiv.1803.04585.
- 47 Tasfia Mashiat, Patrick J. Fowler, and Sanmay Das. Who pays the rent? implications of spatial inequality for prediction-based allocation policies. *CoRR*, abs/2508.08573(2):1686–1697, 2025. doi:10.48550/arXiv.2508.08573.
- 48 Frank D McSherry. Privacy integrated queries: an extensible platform for privacy-preserving data analysis. In *Proceedings of the 2009 ACM SIGMOD International Conference on Management of data*, pages 19–30, 2009. doi:10.1145/1559845.1559850.
- 49 Hervé Moulin. *Fair division and collective welfare*. MIT Press, 2003.
- 50 Harikrishna Narasimhan and Shivani Agarwal. On the relationship between binary classification, bipartite ranking, and binary class probability estimation. In Christopher J. C. Burges, Léon Bottou, Zoubin Ghahramani, and Kilian Q. Weinberger, editors, *Advances in Neural Information Processing Systems 26: 27th Annual Conference on Neural Information Processing Systems 2013. Proceedings of a meeting held December 5-8, 2013, Lake Tahoe, Nevada, United States*, volume 26, pages 2913–2921, 2013. URL: <https://proceedings.neurips.cc/paper/2013/hash/05311655a15b75fab86956663e1819cd-Abstract.html>.
- 51 Fengshi Niu, Harsha Nori, Brian Quistorff, Rich Caruana, Donald Ngwe, and Aadharsh Kannan. Differentially private estimation of heterogeneous causal effects. In *Conference on Causal Learning and Reasoning*, pages 618–633. PMLR, 2022. URL: <https://proceedings.mlr.press/v177/niu22a.html>.

- 52 Juan C. Perdomo, Tolani Britton, Moritz Hardt, and Rediet Abebe. Difficult Lessons on Social Prediction from Wisconsin Public Schools, September 2023. arXiv:2304.06205 [cs]. doi:10.48550/arXiv.2304.06205.
- 53 Juan Carlos Perdomo. The relative value of prediction in algorithmic decision making. *CoRR*, abs/2312.08511, 2023. doi:10.48550/arXiv.2312.08511.
- 54 Krishna Pillutla, Jalaj Upadhyay, Christopher A. Choquette-Choo, Krishnamurthy Dvijotham, Arun Ganesh, Monika Henzinger, Jonathan Katz, Ryan McKenna, H. Brendan McMahan, Keith Rush, Thomas Steinke, and Abhradeep Thakurta. Correlated noise mechanisms for differentially private learning. *CoRR*, abs/2506.08201, 2025. doi:10.48550/arXiv.2506.08201.
- 55 Shir Raviv. When do citizens resist the use of ai algorithms in public policy? theory and evidence. Available at SSRN: <https://ssrn.com/abstract=4328400> or <http://dx.doi.org/10.2139/ssrn.4328400>, January 2023. URL: <https://ssrn.com/abstract=4328400>.
- 56 Matthew J. Salganik, Ian Lundberg, Alexander T. Kindel, Caitlin E. Ahearn, Khaled Al-Ghoneim, Abdullah Almaatouq, Drew M. Altschul, Jennie E. Brand, Nicole Bohme Carnegie, Ryan James Compton, Debanjan Datta, Thomas Davidson, Anna Filippova, Connor Gilroy, Brian J. Goode, Eaman Jahani, Ridhi Kashyap, Antje Kirchner, Stephen McKay, Allison C. Morgan, Alex Pentland, Kivan Polimis, Louis Raes, Daniel E. Rigobon, Claudia V. Roberts, Diana M. Stanescu, Yoshihiko Suhara, Adaner Usmani, Erik H. Wang, Muna Adem, Abdulla Alhajri, Bedoor AlShebli, Redwane Amin, Ryan B. Amos, Lisa P. Argyle, Livia Baer-Bositis, Moritz Büchi, Bo-Ryehn Chung, William Eggert, Gregory Faletto, Zhilin Fan, Jeremy Freese, Tejomay Gadgil, Josh Gagné, Yue Gao, Andrew Halpern-Manners, Sonia P. Hashim, Sonia Hausen, Guanhua He, Kimberly Higuera, Bernie Hogan, Ilana M. Horwitz, Lisa M. Hummel, Naman Jain, Kun Jin, David Jurgens, Patrick Kaminski, Areg Karapetyan, E. H. Kim, Ben Leizman, Naijia Liu, Malte Möser, Andrew E. Mack, Mayank Mahajan, Noah Mandell, Helge Marahrens, Diana Mercado-Garcia, Viola Mocz, Katariina Mueller-Gastell, Ahmed Musse, Qiankun Niu, William Nowak, Hamidreza Omidvar, Andrew Or, Karen Ouyang, Katy M. Pinto, Ethan Porter, Kristin E. Porter, Crystal Qian, Tamkinat Rauf, Anahit Sargsyan, Thomas Schaffner, Landon Schnabel, Bryan Schonfeld, Ben Sender, Jonathan D. Tang, Emma Tsurkov, Austin Van Loon, Onur Varol, Xiafei Wang, Zhi Wang, Julia Wang, Flora Wang, Samantha Weissman, Kirstie Whitaker, Maria K. Wolters, Wei Lee Woon, James Wu, Catherine Wu, Kengran Yang, Jingwen Yin, Bingyu Zhao, Chenyun Zhu, Jeanne Brooks-Gunn, Barbara E. Engelhardt, Moritz Hardt, Dean Knox, Karen Levy, Arvind Narayanan, Brandon M. Stewart, Duncan J. Watts, and Sara McLanahan. Measuring the predictability of life outcomes with a scientific mass collaboration. *Proceedings of the National Academy of Sciences*, 117(15):8398–8403, April 2020. doi:10.1073/pnas.1915006117.
- 57 Robert J Serfling. Probability inequalities for the sum in sampling without replacement. *The Annals of Statistics*, pages 39–48, 1974.
- 58 Ali Shirali, Rediet Abebe, and Moritz Hardt. Allocation Requires Prediction Only if Inequality Is Low, June 2024. arXiv:2406.13882 [cs, econ]. doi:10.48550/arXiv.2406.13882.
- 59 Dan Sperber, David Premack, Ann J. Premack, Dan Sperber, and Ann J. Premack, editors. *The looping effects of human kinds*. Symposia of the Fyssen foundation. Clarendon, Oxford, 1995. doi:10.1093/acprof:oso/9780198524021.001.0001.
- 60 Thomas Steinke and Jonathan Ullman. Tight lower bounds for differentially private selection. In *2017 IEEE 58th Annual Symposium on Foundations of Computer Science (FOCS)*, pages 552–563. IEEE, 2017. doi:10.1109/FOCS.2017.57.
- 61 Stefan Wager and Susan Athey. Estimation and inference of heterogeneous treatment effects using random forests. *Journal of the American Statistical Association*, 113(523):1228–1242, 2018.
- 62 Angelina Wang, Sayash Kapoor, Solon Barocas, and Arvind Narayanan. Against predictive optimization: On the legitimacy of decision-making algorithms that optimize predictive accuracy. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency, FAccT 2023, Chicago, IL, USA, June 12-15, 2023*, volume 1, page 626. ACM, 2023. doi:10.1145/3593013.3594030.

- 63 Boya Wang, Wouter Lueks, Justinas Sukaitis, Vincent Graf Narbel, and Carmela Troncoso. Not Yet Another Digital ID: Privacy-Preserving Humanitarian Aid Distribution. In *2023 IEEE Symposium on Security and Privacy (SP)*, pages 645–663, May 2023. ISSN: 2375-1207. doi:10.1109/SP46215.2023.10179306.
- 64 Bryan Wilder and Pim Welle. Learning treatment effects while treating those in need. In *Proceedings of the 26th ACM Conference on Economics and Computation*, pages 448–473, 2025. doi:10.1145/3736252.3742570.
- 65 World Bank. Revisiting targeting in social assistance: A new look at old challenges. Policy research report, World Bank, 2022. URL: <https://documents1.worldbank.org/curated/en/099300006162240907/pdf/P1648760df4c480270bd0b02288943f413e.pdf>.

A Differences between Our Model and that of Shirali et al. [58]

In their definition of ULA, the authors of [58] make the additional assumption that unit-level administrators can be relied on to allocate aid *within* each unit efficiently, e.g. by avoiding the top q fraction of individuals by welfare with probability $1 - q'$.

We regard this assumption as unsatisfying for two main reasons. Firstly, the allocation scheme described essentially corresponds to a hybrid of ILA and ULA, which could be called “ILA within units.” However, the treatment of ILA within units and full ILA is asymmetric – their model assumes that ILA is forced to expend some of its budget to guarantee accurate targeting across the population, but there is no explanation or cost associated with the accuracy of within-unit targeting. We find this asymmetry unsatisfying. Secondly, in our specific context of *private* allocations, we require the entire process by which allocation decisions are made to satisfy DP. It is therefore important that any decision-making power exercised by unit-level administrators be modeled explicitly to allow for rigorous privacy analysis. To resolve these issues, our model assumes that ULA allocates aid uniformly to all individuals within each unit, without targeting (see Section 2.3.2).

We recognize that the use of within-unit targeting in [58] was partly motivated by the idea that allocating aid to every member of a unit is clearly inefficient in certain contexts. For instance, when units correspond to schools, it may be reasonable to assume that administrators won’t allocate limited tutoring resources to honor-roll students; in this context, within-unit targeting typically wouldn’t raise additional privacy concerns because the identities of honor-roll students are often publicly available to begin with.

We believe that the cleanest way to reconcile this tension and make our model of aid allocation compatible with DP is to assume that targeting within units is uniform while defining units at a finer level of granularity (e.g., “honor-roll students at School A,” “non-honor-roll students at School B,” etc.), a strategy we examine in greater detail in Section 5. More generally, we assume that the population under consideration for both ILA and ULA doesn’t include individuals who are *a priori* ineligible for a benefit, so that resources are not expended on estimating the welfare of irrelevant units.

B Useful Technical Results

The following two lemmas provide high-probability upper bounds on the maximum of Gaussian and hypergeometric random variables, which we utilize in our later proofs in the Appendix.

► **Lemma 29** (Concentration of maximum of Gaussians). *Let $z_1, \dots, z_n$ be (not-necessarily independent) Gaussian random variables with $z_i \sim \mathcal{N}(0, \sigma_i^2)$, and let $\sigma_*^2 = \max_i \sigma_i^2$. Then, with probability at least $1 - \beta$, we have:*

$$\max_i z_i \leq \sigma_* \left(\sqrt{2 \log n} + \sqrt{2 \log(1/\beta)} \right) \quad (29)$$

Proof. Let $Z = \max_i z_i$. Using the standard upper-bound based on moment-generating functions, we have that $\mathbb{E}[Z] \leq \sigma_* \sqrt{2 \log n}$. Then, by the Borell-TIS inequality, we know that the maximum of Gaussians concentrates closely around its expectation [2]:

$$\mathbb{P}(Z > \mathbb{E}[Z] + u) \leq \exp\left(\frac{-u^2}{2\sigma_*^2}\right) \quad (30)$$

Combining these two results and rearranging yields the lemma. $\blacktriangleleft$

► **Lemma 30** (Concentration of hypergeometric random variables ([57] Corollary 1.1)). *Let X be a Hypergeometric random variable with population size P , sample size k , and mean kp . Then with probability at least $1 - \beta$,*

$$|X/k - p| \leq \sqrt{\frac{P-k}{P}} \sqrt{\frac{\log(2/\beta)}{k}} \quad (31)$$

A fundamental tool for achieving zCDP is the *Gaussian mechanism*, which we use as a building block in our algorithms.

► **Lemma 31** (ψ -zCDP Guarantee of the Gaussian Mechanism ([15] Proposition 16)). *Let $q : \mathcal{X}^n \rightarrow \mathbb{R}$ be a sensitivity Δ query. Then the mechanism $M : \mathcal{X}^n \rightarrow \mathbb{R}$ defined by $M(x) \sim \mathcal{N}(q(x), \Delta^2/(2\psi))$ satisfies ψ -zCDP.*

In Section 5, we use the following result of Feldman et al. [21] to characterize the sample complexity of DP-SCO:

► **Theorem 32** (Sample Complexity of DP-SCO ([21] Theorem 4.4)). *Let $\mathcal{X} \times [0, 1]$ be the domain of datasets, and $\mathcal{D}$ be a distribution over $\mathcal{X} \times \{0, 1\}$. Let $S = ((x_1, y_1) \dots, (x_n, y_n))$ be a dataset drawn i.i.d. from $\mathcal{D}$. Let $\mathcal{K} \subseteq \mathbb{R}^P$ be a convex set denoting the space of all models. Let $\ell : \mathcal{K} \times \mathcal{X} \times \{0, 1\} \rightarrow \mathbb{R}$ be a loss function, which is L -Lipschitz, ξ -smooth, and convex in its first parameter. Then, if $\mathcal{K}$ has diameter at most D , the output $\hat{\theta}$ of Algorithm 2 in [21] satisfies ψ -zCDP and we have:*

$$\mathbb{E}[\mathcal{L}(\hat{\theta})] \leq \min_{\theta^* \in \mathcal{K}} \mathcal{L}(\theta^*) + 10LD \left(\frac{1}{\sqrt{n}} + \frac{\sqrt{p}}{n\sqrt{2\psi}} \right) \quad (32)$$

where $\mathcal{L}(\theta) = \mathbb{E}_{x, y \sim \mathcal{P}}[\ell(\theta, x, y)]$, provided that $\frac{D}{L} \min \left\{ \frac{4}{\sqrt{n}}, \frac{\sqrt{2\psi}}{\sqrt{p}} \right\} \leq 2/\xi$.

C Missing Proofs

In the following, we provide proofs of the key results from Section 2 concerning the utility of ILA, ULA, and RAND. Proofs of results from the later sections can be found online in the full version of the paper.¹

Proof of Lemma 4. Let $k' \leq k$ be given. As in the lemma statement, let $\mathcal{I}$ denote the set of individuals mistakenly included or excluded from our k' -allocation because of estimation error. Maximizing the normalized expected value of ILA can be represented as a linear program:

$$\max_x \sum_{i=1}^P x_i(1 - \eta_i) \quad s.t. \quad \sum_{i=1}^P x_i = k' \quad (33)$$

$$0 \leq x_i \leq 1 \quad \forall i \quad (34)$$

¹ <https://arxiv.org/abs/2604.15596>

Since the optimal solution for an integer k' is to assign $x_{s(i)} = 1$ for all $i \leq k'$, we do not need to worry about non-integral solutions.

Let x^* denote the optimal solution for the true η vector, and $\tilde{x}^*$ be the optimal solution for the perturbed $\hat{\eta}$ induced by our estimated probabilities. Then, letting $\langle \cdot, \cdot \rangle$ denote the standard inner product, we have:

$$\langle \tilde{x}^*, 1 - \eta \rangle = \langle \tilde{x}^*, 1 - \hat{\eta} \rangle - \langle \tilde{x}^*, \eta - \hat{\eta} \rangle \quad (35)$$

$$\geq \langle x^*, 1 - \hat{\eta} \rangle - \langle \tilde{x}^*, \eta - \hat{\eta} \rangle \quad (36)$$

$$= \langle x^*, 1 - \eta \rangle - \langle x^*, \hat{\eta} - \eta \rangle - \langle \tilde{x}^*, \eta - \hat{\eta} \rangle \quad (37)$$

$$= \langle x^*, 1 - \eta \rangle - \langle x^* - \tilde{x}^*, \hat{\eta} - \eta \rangle \quad (38)$$

$$= \langle x^*, 1 - \eta \rangle - \sum_{i=1}^P (x_i^* - \tilde{x}_i^*) (\hat{\eta}_i - \eta_i) \quad (39)$$

Subtracting $\langle \tilde{x}^*, 1 - \eta \rangle$ and adding $\sum_{i=1}^P (x_i^* - \tilde{x}_i^*) (\hat{\eta}_i - \eta_i)$ to both sides gives us that our regret is at most:

$$\sum_{i=1}^P (x_i^* - \tilde{x}_i^*) (\eta_i - \hat{\eta}_i) \leq \sum_{i \in \mathcal{I}} |\eta_i - \hat{\eta}_i| \quad (40)$$

From here, we have that in expectation, the regret of x^* with respect to the optimal k' -allocation in hindsight (i.e. based on the realized welfares) is exactly $\sum_{i=1}^P x_i^* \eta_i = \sum_{i=1}^{k'} \eta_{s(i)} =: \bar{\eta}_{k'}$.

Finally, the normalized regret with respect to the optimal k -allocation is then at most that quantity plus $(k - k')$. $\blacktriangleleft$

Proof of Lemma 6. In Shirali et al. [58] (Lemma 4.1), the authors prove that the value of ULA is lower-bounded by $\sum_{j=1}^{k/N} N_{s(j)} (1 - \rho_{s(j)})$, implying that $\text{REGRET}(\text{ULA}, w; k) \leq \sum_{j=1}^{k/N} N_{s(j)} \rho_{s(j)}$. An immediate generalization says that if we allocate aid to $x_j \in [0, N_j]$ individuals chosen uniformly at random in unit j , then our normalized regret is upper-bounded by $\sum_{j=1}^M x_j \rho_j$. This allows us to formulate ULA as a linear program:

$$\min_x \sum_{j=1}^M x_j \rho_j \quad \text{s.t.} \quad \sum_{j=1}^M x_j = k \quad (41)$$

$$0 \leq x_j \leq N_j \quad \forall j \quad (42)$$

Let x^* denote the optimal solution for the true unit profile vector ρ , and let $\tilde{x}^*$ denote the optimal solution for $\tilde{\rho}$. Then:

$$\langle \tilde{x}^*, \rho \rangle = \langle \tilde{x}^*, \tilde{\rho} \rangle + \langle \tilde{x}^*, \rho - \tilde{\rho} \rangle \quad (43)$$

$$\leq \langle x^*, \tilde{\rho} \rangle + \langle \tilde{x}^*, \rho - \tilde{\rho} \rangle \quad (44)$$

$$= \langle x^*, \rho \rangle + \langle x^*, \tilde{\rho} - \rho \rangle + \langle \tilde{x}^*, \rho - \tilde{\rho} \rangle \quad (45)$$

$$= \langle x^*, \rho \rangle + \langle x^* - \tilde{x}^*, \tilde{\rho} - \rho \rangle \quad (46)$$

$$= \langle x^*, \rho \rangle + \sum_{j=1}^M (x_j^* - \tilde{x}_j^*) (\tilde{\rho}_j - \rho_j) \quad (47)$$

$$\leq \langle x^*, \rho \rangle + \|x^* - \tilde{x}^*\|_p \|\tilde{\rho} - \rho\|_q \quad (48)$$

where the first inequality follows from the optimality of $\langle \tilde{x}^*, \tilde{\rho} \rangle$ and the last line holds for any $p, q \in [1, \infty]$ such that $1/p + 1/q = 1$ by Hölder's inequality. The last step to derive the general bound is to note that on a worst-case input where all individuals have welfare $w \in \{0, 1\}$, the inequality $\langle x^*, \rho \rangle \geq \overline{\text{REGRET}}(\text{ULA}, w, k)$ is in fact tight.

For the specific choices of $p, q \in \{1, 2, \infty\}$, we argue as follows. Let $\vec{N} = (N, N, \dots, N) \in \mathbb{R}^M$. By the triangle inequality, for any $p \in [1, \infty]$, we have the two following upper bounds:

$$\|x^* - \tilde{x}^*\|_p \leq \|x^*\|_p + \|\tilde{x}^*\|_p \quad (49)$$

$$\|x^* - \tilde{x}^*\|_p = \|(x^* - \vec{N}) - (\tilde{x}^* - \vec{N})\|_p \leq \|x^* - \vec{N}\|_p + \|\tilde{x}^* - \vec{N}\|_p \quad (50)$$

And so in particular, we can always choose the smaller of the two upper bounds. For $p = 1$, we have $\|x^*\|_1 = k$ and $\|x^* - \vec{N}\|_1 = P - k$, giving a final symmetric upper bound of $2 \min(k, P - k)$. For $p = 2$, we have $\|x^*\|_2 \leq \sqrt{kN}$ because x^* has at most k/N components with squared magnitude N^2 , and symmetrically $\|x^* - \vec{N}\|_2 \leq \sqrt{(P - k)N}$. Combining these results yields the desired bounds:

$$\begin{cases} 2 \min(k, P - k) \max_j |\tilde{\rho}_j - \rho_j| & p = 1, q = \infty \\ \sqrt{2 \min(k, P - k) N \sum_{j=1}^M (\tilde{\rho}_j - \rho_j)^2} & p = 2, q = 2 \\ N \sum_{j=1}^M |\tilde{\rho}_j - \rho_j| & p = \infty, q = 1 \end{cases} \quad (51)$$

◀

Proof of Lemma 8. Let X denote the number of low-welfare individuals selected by RAND. Then $\text{VAL}(\text{RAND}) \geq \delta_w X$, which implies that $\max_{|\mathcal{J}|=k} \text{VAL}(\mathcal{J}, w) - \text{VAL}(\text{RAND}, w) \leq \max_{|\mathcal{J}|=k} \text{VAL}(\mathcal{J}, w) - \delta_w X \leq \delta_w k - \delta_w X$, and therefore that $\overline{\text{REGRET}}(\text{RAND}, w; k) \leq k - X$. Observing that X follows a hypergeometric distribution completes the proof. ◀

Proof of Lemma 9. The proof proceeds by maximizing G subject to equality constraints on $\bar{\rho}_M$ and $\bar{\rho}_K$. For simplicity, we ignore the constraint that $\rho_i \in [0, 1]$; this can only increase the value of the objective function and so our final bound remains valid. Without this constraint, it is clear by inspection of the second representation of the Gini coefficient, $G \propto \sum_{1 \leq i \leq j \leq M} (\rho_{s(j)} - \rho_{s(i)})$, that it is maximized when $M - 1$ units have the same welfare and a single unit has higher welfare. So, we set:

$$\rho_{s(i)} = a \quad \forall i \leq M - 1 \quad (52)$$

$$\rho_{s(M)} = a + b \quad (53)$$

$$\bar{\rho}_M = a + \frac{b}{M} \quad (54)$$

$$\bar{\rho}_K = a \quad (55)$$

$$\Delta_K := \bar{\rho}_M - \bar{\rho}_K = \frac{b}{M} \quad (56)$$

Then, using the definition of the Gini coefficient, we have:

$$G = \frac{1}{M^2 \bar{\rho}_M} (M - 1)b = \frac{1}{M^2 \bar{\rho}_M} (M - 1) \frac{Mb}{M} = \frac{1}{M \bar{\rho}_M} (M - 1) \Delta_K \quad (57)$$

$$\Rightarrow \Delta_K = \frac{GM \bar{\rho}_M}{M - 1} > G \bar{\rho}_M \quad (58)$$

as desired. ◀

C.1 PrivatePartialSums

There is an extensive body of work on the problem of privately computing partial sums of bounded data points, which is described in detail in [54]. We use **PrivatePartialSums** to refer to the factorization-based algorithm studied in [28], whose utility guarantees are summarized in [54, Theorem 2.7]. Briefly, their algorithm is based on adding correlated Gaussian noise to each of the n partial sums of the input sequence with maximum standard deviation:

$$\frac{1}{\sqrt{\psi}} \left(1 + \frac{\log(n) + \gamma_E}{\pi} \right) \tag{59}$$

where $\gamma_E \approx 0.577$ denotes Euler's constant, which, when combined with Lemma 29, yields the error term in Theorem 12.