

Avicaching: A Two Stage Game for Bias Reduction in Citizen Science

Yexiang Xue
Computer Science Dept
Cornell University
yexiang@cs.cornell.edu

Ian Davies, Daniel Fink,
Christopher Wood
Cornell Lab of Ornithology
{id99, daniel.fink,
chris.wood}@cornell.edu

Carla P. Gomes
Computer Science Dept
Cornell University
gomes@cs.cornell.edu

ABSTRACT

Citizen science projects have been very successful at collecting rich datasets for different applications. However, the data collected by the citizen scientists are often biased, more aligned with the citizens' preferences rather than scientific objectives. We introduce a novel two-stage game for reducing data-bias in citizen science in which the game *organizer*, a citizen-science program, incentivizes the *agents*, the citizen scientists, to visit under-sampled areas. We provide a novel way of encoding this two-stage game as a single optimization problem, cleverly folding (an approximation of) the agents' problems into the organizer's problem. We present several new algorithms to solve this optimization problem as well as a new structural SVM approach to learn the parameters that capture the agents' behaviors, under different incentive schemes. We apply our methodology to *eBird*, a well-established citizen-science program for collecting bird observations, as a game called *Avicaching*. We deployed Avicaching in two New York counties (March 2015), with a great response from the birding community, surpassing the expectations of the *eBird* organizers and bird-conservation experts. The field results show that the Avicaching incentives are remarkably effective at encouraging the bird watchers to explore under-sampled areas and hence alleviate the *eBird*'s data bias problem.

Keywords

Two-Stage Game, Bilevel Optimization, Structural SVM, Citizen Science

1. INTRODUCTION

Over the past decade, along with the emergence of the *big data* era, the data collection process for scientific discovery has evolved dramatically. One effective way of collecting large datasets is to engage the public through citizen science projects, such as *Zooniverse*, *Cicada Hunt* and *eBird* [24, 42, 35]. The success of these projects relies on the ability to tap into the intrinsic motivations of the volunteers to make participation enjoyable [5]. Thus in order to engage large groups of participants, citizen science projects often have few restrictions, leaving many decisions about

Appears in: *Proceedings of the 15th International Conference on Autonomous Agents and Multiagent Systems (AAMAS 2016)*, J. Thangarajah, K. Tuyls, C. Jonker, S. Marsella (eds.), May 9–13, 2016, Singapore.
Copyright © 2016, International Foundation for Autonomous Agents and Multiagent Systems (www.ifaamas.org). All rights reserved.

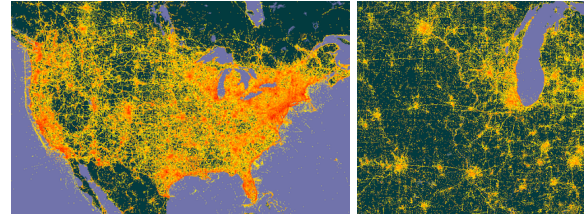


Figure 1: Highly biased distribution of *eBird* observations until 2014. (Left) continental US (Right) Zoom in Midwest US. Submissions coincide with urban areas.

where, when, and how to collect data up to the participants. As a result, the data collected by volunteers are often biased, more aligned with their preferences, rather than providing systematic observations across various experimental settings. Moreover, since participants volunteer their effort, personal convenience is an important factor that often determines how data are collected. For spatial data, this means more searches occur in areas close to urban areas and roads (Fig. 1).

We provide a general methodology to mitigate the data bias problem, as a two-stage game in which the game organizer, e.g., a citizen-science program, provides incentives to the agents, the citizen scientists, to perform more crucial scientific tasks. We apply it to *eBird*, a well-established citizen-science program for collecting bird observations, as a game called *Avicaching*.

Our proposed two-stage game is related to the Principal-Agent framework, originally studied in economics [31], and more recently in computer science [1, 17, 14], and to the Stackelberg games [13, 12, 28, 15], which also involves e.g., a *principal* or a *leader* and *agents* or *followers*. These games have been studied under different assumptions regarding the agents' preferences and computational abilities [18, 8]. In crowdsourcing, there has been related work on mechanisms to improve the crowd performance [29, 23, 21, 22, 34, 26, 37, 7]. Notable works include using incentives to promote exploration activities [16], and steering user participation with badges [3]. [32, 10, 9] discuss the optimal reward allocation to reduce the empirical risk of machine learning models.

In our two-stage game setting, the *agents* are citizen scientists maximizing their intrinsic utilities, as well as the incentives distributed by the game organizer, subject to a budget constraint. The organizer corresponds to an organization with notable influence on the citizen scientists. The

organizer factors in the reasoning process of the citizen scientists to optimize an incentive scheme. In our setting, the game organizer’s goal is to optimize an incentive scheme in order to **induce a uniform data collection process**. Furthermore, the organizer **explicitly models the discrete choice problem of each agent as a knapsack problem**. We refer to this two-stage game as the **Optimal Incentives for Knapsack Agents (OptIKA)** game.

We provide several novel algorithms to solve **OptIKA** and in particular we convert the two-stage game into a single optimization problem by cleverly folding (an approximation of) the agents’ problems into the organizer’s problem.

We consider **(1) different objectives** for the organizer, corresponding to different measures of data uniformity using *Mixed Integer Programming* and *Mixed Integer Quadratic Programming* formulations. We also consider **(2) different levels of rationality** for the *agents*, which result in different approaches to fold the agents’ knapsack problems into the organizer’s problem. For the scenario in which the agents have **unbounded rationality**, we developed an **iterative, row generation method**, given the exponential number of constraints induced by agents’ knapsack problems. We also consider two scenarios in which the agents have **bounded rationality**: one in which the agents use a **greedy heuristic** and another one based on a **dynamic programming (DP)**, **polynomial time approximation scheme** for the knapsack problem. For **(3) scalability**, we use the Taylor expansion of the L2-norm and develop a novel approach based on the *Alternating Direction Method of Multipliers*. **(4)** We propose a **novel structural SVM** framework to solve the so-called **identification problem**, which learns agents’ behaviors under different incentive schemes.

We applied our methodology to *eBird* as a game called **Avicaching**, **deploying it as a pilot study in two New York counties**. Since the inception in March 2015, **19%** of the *eBird* observations in our pilot counties shifted from traditional locations to *Avicaching* locations with no previous observations. Our field results show that our **Avicaching incentives are remarkably effective at encouraging the bird watchers to explore under-sampled areas and hence it alleviates the data bias problem in *eBird***. We also showed that **under our Avicaching scheme, agents can cover the area more uniformly**, which leads to **higher performance on a predictive model for bird occurrence** than the no-incentive case, with the same amount of effort devoted. Our methodology is general and can be applied to other citizen science applications as well as similar scenarios, beyond citizen science.

2. PROBLEM FORMULATION

We consider the setting in which citizen scientists are encouraged to conduct *scientific surveys*. For example in *eBird*, bird watchers survey a given area, and record all the interesting species observed in that area. This setting can be generalized to other scientific exploration activities. The general formulation of the two-stage game is:

$$\begin{aligned}
 & \textbf{(Organizer)} \quad \text{maximize}_r \quad U_o(v, r), \\
 & \quad \text{subject to} \quad B_o(r), \\
 & \textbf{(Agents)} \quad \text{maximize}_v \quad U_a(v, r), \\
 & \quad \text{subject to} \quad B_a(v),
 \end{aligned} \tag{1}$$

where r is the external reward that the organizer (e.g., a cit-

izen science program) uses to steer the agents (e.g., citizen scientists), and v are the reactions from the agents, which is the result of optimizing their own utilities. $U_o(v, r)$ and $U_a(v, r)$ are the utility functions of the organizer and agents, respectively, and $B_o(r)$ and $B_a(v)$ are their respective budget constraints.

The Organizer’s Objective is to promote a balanced exploration activity, which corresponds to sending people to under-sampled areas. **The pricing problem** is the associated organizer’s problem of determining the optimal rewards to induce the desired behavior from the agents, namely sending the agents to undersampled areas. Let $L = \{l_1, l_2, \dots, l_n\}$ be the set of locations, and $X_{0,i}$ the number of historical visits at location l_i at the beginning of a time period T . Suppose there are m citizen scientists $b_1, b_2, \dots, b_m$. During time period T , each citizen scientist b_j chooses a set $L_j \subseteq L$ of locations to explore. At the end of time period T , location i received a net amount of visits $V_i = |\{l_i \in L_j : j = 1, \dots, m\}|$ and its total number Y_i of visits corresponds to $Y_i = X_{0,i} + V_i$. We denote by $\mathbf{Y}$ the column vector $(Y_1, \dots, Y_n)^T$ and by $\bar{\mathbf{Y}}$ the constant column vector $(\bar{Y}, \dots, \bar{Y})^T$ where $\bar{Y} = \frac{1}{n} \sum_{i=1}^n Y_i$. As the organizer aims to promote a more uniform sampling effort among different locations, this objective can be expressed as the reduction $D_p = \frac{1}{n} \|\mathbf{Y} - \bar{\mathbf{Y}}\|_p^p$ of the difference between $\mathbf{Y}$ and $\bar{\mathbf{Y}}$. Given this definition, D_1 corresponds to the *mean absolute deviation*, while D_2 corresponds to the *sample variance*. Other objectives could be used, e.g., maximizing the entropy of the sample distribution in order to minimize the distance to a uniform distribution.¹ (See section 3.1.)

The Agents’ Model – Each agent is maximizing her own utility subject to her budget constraint. Namely, if a citizen scientist b_j chooses to visit location l_i , she will receive an intrinsic utility $u_{j,i}$, at a cost $c_{j,i}$. We assume that agent b_j has a given budget C_j , so the total cost of all the places explored by b_j cannot exceed C_j .

To incentivize citizen scientists to visit undersampled areas, the organizer introduces an extra incentive r_i for each location l_i . Every citizen scientist visiting location l_i receives an extra reward r_i , besides their internal utility $u_{j,i}$. For the sake of fairness, we require that these rewards only vary across locations and are the same for all agents. In addition, to make it easier to communicate with the agents, we assume that all rewards come from a fixed discrete set: $r_i \in R = \{r_1^*, \dots, r_k^*\}$. Taking into account intrinsic utilities, external rewards and the budget constraint, the citizen scientist b_j ’s planning problem becomes:

$$\begin{aligned}
 & \text{maximize}_{L_j \subseteq L} \quad \sum_{l_i \in L_j} u_{j,i} + w_r \cdot r_i, \\
 & \text{subject to} \quad \sum_{l_i \in L_j} c_{j,i} \leq C_j.
 \end{aligned} \tag{2}$$

In this formulation, $u_{j,i}$ is the intrinsic utility, r_i is the external reward, w_r is the relative importance ratio between the intrinsic utilities and the external rewards, $c_{j,i}$ is the cost, and C_j is the total budget for an agent. Overall, combining the organizer’s goal and the agents’ models, the pricing

¹Note: uncertainty measures, often used in active learning [30], are typically tied to a particular predictive model and therefore do not serve our goal of meeting multiple scientific objectives with balanced sampling. We cannot commit to improving one particular predictive model.

Organizer's Objective	Agents' Rationality			
	Bounded			
	Full	DP	Greedy	
L1-Norm	OptIKA-L1-Full	OptIKA-L1-DP	OptIKA-L1-Greedy	MIP
L2-Norm	OptIKA-L2-Full	OptIKA-L2-DP	OptIKA-L2-Greedy	MIQP
L2-Taylor	OptIKA-L2T-Full	OptIKA-L2T-DP	OptIKA-L2T-Greedy	MIP/ADMM
	Iterative Row Gen			Single Constraint Programming instance

Figure 2: Two stage game scenarios and corresponding algorithms described in Section 3.1.

problem of the **Optimal Incentives for Knapsack Agents (OptIKA)** game is:

$$\begin{aligned}
(\text{OptIKA}) \quad & \underset{\mathbf{r}}{\text{minimize}} \quad \frac{1}{n} \|\mathbf{Y} - \bar{\mathbf{Y}}\|_p^p \\
& \text{subject to} \quad L_j = \underset{L_j \subseteq L}{\text{argmax}} \sum_{i \in L_j} u_{j,i} + w_r \cdot r_i, \\
& \quad \sum_{l_i \in L_j} c_{j,i} \leq C_j, \quad j \in \{1, \dots, m\}, \\
& \quad r_i \in R, \quad i \in \{1, \dots, n\}.
\end{aligned} \tag{3}$$

The **Identification Problem** learns the parameters capturing agents' behavior, by fitting a model to predict the agents' preferences under various incentive schemes. The identification problem is related to Inverse Reinforcement Learning [27, 36, 20], in which one also assumes that the agents are optimizing for long-term rewards (Section 3.2).

3. ALGORITHMS

3.1 Pricing Problem

We developed a variety of algorithms to solve the pricing problem, capturing different organizer's objectives and agents' computational capabilities, as summarized in Fig. 2. First, the complexity of the **OptIKA** problem depends on the uniformity measure of the organizer. The **OptIKA** problem can be solved using a Mixed Integer Programming (MIP) formulation when minimizing the mean absolute deviation (L_1 -norm), whereas it becomes a Mixed Integer Quadratic Program (MIQP) when minimizing the sample variance (L_2 -norm). Second, the computational capability of the agents impacts how one can fold the constraints of the agents into the organizer's problem. If the agent solves her knapsack problem optimally with *full* rationality, it yields an exponential number of constraints to be handled by the organizer, thus raising scalability issues and requiring an iterative approach (see *Row Generation* encoding). We also consider the case in which agents have *bounded* rationality, whether the agent solves her knapsack using a *dynamic programming*-based approach or in a *greedy* fashion. In this scenario, the polynomial number of linear constraints to consider can be encoded in a single Constraint Programming instance. Furthermore, in order to scale up with the number of agents, we improve our approach with a decomposition method that decouples the agents' optimization problems (see *ADMM*).

3.1.1 Modeling the Organizer's Objective

A first measure of sample uniformity is the mean absolute deviation D_1 , which allows us to formulate the objective function as a MIP. For every location l_i , introduce a variable Z_i such that $Z_i \geq |Y_i - \bar{Y}|$. Overall, the organizer's objec-

Algorithm 1: Row Generation OptIKA-LX-Full

```

1  $\Phi \leftarrow \emptyset$ ;
2  $OptimalFlag \leftarrow False$ ;
3 while  $OptimalFlag = False$  do
4    $(\mathbf{r}_\dagger, \mathbf{v}_{1\dagger}, \dots, \mathbf{v}_{m\dagger}) \leftarrow \text{OptIKA-LX-Full-Relax}(\Phi)$ ;
5    $OptimalFlag \leftarrow True$ ;
6   for  $j \in \{1, \dots, m\}$  do
7      $\mathbf{v}_j^* \leftarrow \underset{\text{subject to } \mathbf{c}_j^T \cdot \mathbf{v}_j \leq C_j}{\text{argmax}} (\mathbf{u}_j + w_r \cdot \mathbf{r}_\dagger)^T \cdot \mathbf{v}_j$ ;
8     if  $(\mathbf{u}_j + w_r \cdot \mathbf{r}_\dagger)^T \cdot \mathbf{v}_j^* > (\mathbf{u}_j + w_r \cdot \mathbf{r}_\dagger)^T \cdot \mathbf{v}_{j\dagger}$  then
9        $\Phi \leftarrow \Phi \cup \{ (\mathbf{u}_j + w_r \cdot \mathbf{r})^T \cdot \mathbf{v}_j \geq$ 
10          $(\mathbf{u}_j + w_r \cdot \mathbf{r})^T \cdot \mathbf{v}_j^* \}$ ;
11        $OptimalFlag = False$ ;
12   end
13 end
14 end

```

tive function can be captured as: $\min \sum_{i=1}^n Z_i$, s.t. $Z_i \geq Y_i - \bar{Y}$, $Z_i \geq \bar{Y} - Y_i$. We refer to this formulation as **OptIKA-L1**. A second formulation (**OptIKA-L2**) uses the L_2 -norm sample variance (D_2). In this case, the organizer's objective is quadratic, and hence the entire problem becomes a Mixed Integer Quadratic Program (MIQP). As a third option, we model the organizer's objective using the Taylor approximation of the sample variance (**OptIKA-L2T**), in which case the organizer's problem translates into minimizing $S = \sum_{i=1}^n s_i V_i$, where $s_i = \frac{2}{n} (X_{0,i} - \bar{X}_0)$. Notice that the Stackelberg pricing games studied in [17] is a special case of **OptIKA** in this form, therefore **OptIKA** is APX-hard. We leave the details in the supplementary materials, anonymously online at [40].

3.1.2 Modeling the Knapsack Agents

Row Generation Encoding We first present the algorithm **OptIKA-LX-Full** (where X is either 1, 2 or 2T) in which we assume the agents have full rationality, and their reasoning process is captured by an iterative row generation method. This algorithm can be combined with any of the three organizer's objectives. Let $v_{j,i}$ be a binary variable, which is 1 if and only if $l_i \in L_j$. Using vector representations, we set $\mathbf{v}_j = (v_{j,1}, \dots, v_{j,n})^T$, $\mathbf{u}_j = (u_{j,1}, \dots, u_{j,n})^T$, $\mathbf{c}_j = (c_{j,1}, c_{j,2}, \dots, c_{j,n})^T$, $\mathbf{r} = (r_1, \dots, r_n)^T$, $\mathbf{s} = (s_1, \dots, s_n)^T$.

The key to solve a bilevel optimization like **OptIKA** is to find clever ways to fold the nested optimization (the agents' problem) as constraints to the global problem. Here we show that the agents' optimization problem can be transformed into an exponential number of constraints of the type:

$$(\mathbf{u}_j + w_r \cdot \mathbf{r})^T \cdot \mathbf{v}_j \geq (\mathbf{u}_j + w_r \cdot \mathbf{r})^T \cdot \mathbf{v}_j', \tag{4}$$

in which $\mathbf{v}_j'$ ranges over all vectors in $\{0, 1\}^n$, which satisfies $\mathbf{c}_j^T \cdot \mathbf{v}_j' \leq C_j$, for all $j \in \{1, \dots, m\}$. The intuitive meaning of Inequality 4 is that the location set that the agent chooses is better in terms of utility values than any other location set within the budget constraint. We use Φ to denote a set of constraints of this form, and we write $\{\mathbf{v}_1, \dots, \mathbf{v}_m\} \in \Phi$ to mean that $\mathbf{v}_1, \dots, \mathbf{v}_m$ satisfy all the constraints in Φ .

We cannot add all the constraints upfront, as there are exponentially many of them. Instead, we add them in an iterative manner until proving optimality. The row generation scheme starts by solving a relaxation of the original pric-

ing problem: **OptIKA-LX-Full-Relax**(Φ), with a small initial constraint set Φ of constraints as shown in Inequality 4:

$$\begin{aligned} \text{OptIKA-LX-Full-Relax}(\Phi) : \text{Min: (organizer's obj)} \quad & D_p \text{ or } S, \\ \text{s. t. } \quad & \mathbf{c}_j^T \cdot \mathbf{v}_j \leq C_j, \quad j \in \{1, \dots, m\}, \\ & \{\mathbf{v}_1, \dots, \mathbf{v}_m\} \in \Phi, \\ & r_i \in R, \quad i \in \{1, \dots, n\}. \end{aligned}$$

Then the algorithm seeks to enlarge the set Φ with new constraints of the form in Inequality 4 to further improve the objective function. This step is done by solving the Knapsack problem for each agent. If the current response of one agent b_j is not the optimal response to the Knapsack problem, then it implies that at least one constraint of the form shown in Inequality 4 is violated. We then add in the constraint into Φ and solve again. The whole algorithm iterates until no new constraints can be added to Φ , at which point we can prove optimality. The algorithm is shown as Algorithm 1.

There is one last subtlety: the Inequality (4) is not a linear one, because both $\mathbf{r}$ and $\mathbf{v}_j$ are variables. To linearize it, we bring in an extra variable $ur_{j,i}$, and we add constraints to ensure that $ur_{j,i}$ is always equal to $v_{j,i} \cdot (u_{j,i} + w_r \cdot r_i)$. The constraints needed are:

$$\begin{aligned} ur_{j,i} &\geq 0, \\ ur_{j,i} &\leq u_{j,i} + w_r \cdot r_i, \\ v_{j,i} = 0 &\Rightarrow ur_{j,i} \leq 0, \end{aligned} \quad (5)$$

In this case, Inequality (4) can be rewritten as $\sum_{i=1}^n ur_{j,i} \geq (\mathbf{u}_j + w_r \cdot \mathbf{r})^T \cdot \mathbf{v}_j$. Eq. (5) is an indicator constraint, which can be linearized with the big-M formulation [11].

Dynamic Programming Encoding **OptIKA-LX-DP** is the best performing encoding, motivated by the polynomial time approximation scheme to solve the Knapsack Problem. It encodes the entire problem into a single MIP, rather than a series of iterative MIPs as in the row generation approach. It reflects the bounded rationality of agents, as it sacrifices a little precision when modeling the cost under a bounded memory size.

In the Knapsack Problem for citizen scientist b_j , we first discretize the budget C_j into N_b equal-sized units. Let $\mathcal{D}_j = \{kC_j/N_b | k = 0, \dots, N_b\}$ be the set of all discrete units. We further round the cost $c_{j,i}$ to its nearest discrete unit from above in $\mathcal{D}_j$. We introduce extra variables $opt(j, i, c)$, for $i \in \{1, \dots, n\}$ and $c \in \mathcal{D}_j$, to denote the optimal utility for agent b_j if we only consider the first i locations $l_1, \dots, l_i$ and the total cost cannot exceed c . Consider the Dynamic Programming recursion to solve the Knapsack Problem:

$$opt(j, i, c) = \begin{cases} \max\{opt(j, i-1, c - c_{j,i}) + u_{j,i} + w_r \cdot r_i, \\ \quad \quad \quad opt(j, i-1, c)\}, & \text{if } i > 1, c \geq c_{j,i}, \\ opt(j, i-1, c), & \text{if } i > 1, c < c_{j,i}, \\ 0, & \text{otherwise.} \end{cases} \quad (6)$$

The key insight of **OptIKA-LX-DP** is that this recursion can be translated as a set of linear inequalities. As an example, when $i > 1$ and $c \geq c_{j,i}$, the recursion can be encoded as,

$$opt(j, i, c) \geq opt(j, i-1, c), \quad (7)$$

$$opt(j, i, c) \geq opt(j, i-1, c - c_{j,i}) + u_{j,i} + w_r \cdot r_i. \quad (8)$$

There are similar inequalities to capture other cases in Equation 6. Denote $u_{\mathcal{D}_j}^{knaps}$ as the optimal utility for solving

the knapsack problem for citizen scientist b_j . We must have $u_{\mathcal{D}_j}^{knaps} \geq opt(j, n, c)$, for all $c \in \mathcal{D}_j$. In summary, agent b_j 's knapsack problem can be encoded as:

$$\begin{aligned} (\mathbf{u}_j + w_r \cdot \mathbf{r})^T \cdot \mathbf{v}_j &\geq u_{\mathcal{D}_j}^{knaps}, \\ \mathbf{c}_j^T \cdot \mathbf{v}_j &\leq C_j \text{ and } u_{\mathcal{D}_j}^{knaps} \geq opt(j, n, c), \quad \forall c \in \mathcal{D}_j, \end{aligned} \quad (9)$$

Here, $opt(j, n, c)$ is encoded by linear inequalities similar to the ones in Equations 7 and 8. Finally, we need to use a big-M notation to linearize the inequality in Equation 9.

If we bound N_b , this encoding introduces $O(mnN_b)$ extra variables and $O(mnN_b)$ extra constraints. Notice that this encoding can be combined with the row generation approach. We can first solve the problem under limited precision using this dynamic programming encoding, then further refine the solution using the row generation approach.

Greedy Formulation **OptIKA-LX-Greedy** assumes each agent follows a simple greedy heuristic: first, rank all the locations based on their efficiency, i.e. the ratio between the utility (including the external reward) and the cost; then greedily select locations with the highest efficiency, without exceeding the budget limit. This simple heuristic is a 2-approximation for the Knapsack problem, and works well in practice. Define $\psi_{j,i} = (w_r \cdot r_i + u_{j,i})/c_{j,i}$ as the efficiency of location l_i according to agent b_j . Our formulation is based on the following theorem:

THEOREM 1. *Assume for all $i \neq i'$, $\psi_{j,i} \neq \psi_{j,i'}^2$, then $\{v_{j,1}, \dots, v_{j,n}\}$ is a decision made by the greedy algorithm if and only if the following two constraints hold:*

$$v_{j,i} = 0 \Rightarrow c_{j,i} > C_j - \sum_{i' \neq i} v_{j,i'} c_{j,i'} \mathbf{1}(\psi_{j,i'} \geq \psi_{j,i}), \quad (10)$$

for all $i \in 1, \dots, n$, and $\sum_{i=1}^n c_{j,i} \cdot v_{j,i} \leq C_j$.

In this theorem, $\mathbf{1}(\psi_{j,i'} \geq \psi_{j,i})$ is an indicator variable, which is one if and only if $\psi_{j,i'} \geq \psi_{j,i}$. Theorem 1 translates the greedy process into a set of constraints. The intuitive meaning of inequality (10) says that if location l_i is not in the knapsack ($v_{j,i} = 0$), then it must be the case that some locations with higher efficiency than l_i has already taken up its space. Unfortunately, inequality (10) is not linear. We can again use big-M notation to transform the above constraint into a set of linear constraints.

3.1.3 Scaling to Many Agents with ADMM

In order to model a large number of citizen scientists, the pricing algorithm needs to be able to scale. To that end, we develop **OptIKA-L2T-ADMM**, harnessing a variant of the Alternating Direction Method of Multipliers [6, 25]. This approach decomposes the global problem of designing the rewards for all agents to a series of subproblems, each of which designs the optimal rewards for one agent. Then the algorithm matches the local rewards for all agents using an iterative approach. To the best of our knowledge, this is the first time that a decomposition based method is introduced to solve the optimal pricing problem. Because ADMM requires a decomposable objective function, this variant only applies to the third organizer's objective function that uses the Taylor expansion (**OptIKA-L2T**) as described in the supplementary material [40]. We introduce a local copy of the

²In practice, efficiencies almost always differ when they are learned from data.

reward vector for each agent b_j : $\mathbf{r}_j = (r_{j,1}, \dots, r_{j,n})^T$, and we rewrite the global problem as:

$$\begin{aligned} \min \quad & S = \sum_{j=1}^m \mathbf{s}^T \cdot \mathbf{v}_j, \\ \text{s.t.} \quad & (\mathbf{r}_j, \mathbf{v}_j) \in \Sigma_j, \quad \mathbf{r}_j = \mathbf{r}, \quad \forall j \in \{1, \dots, m\}. \end{aligned}$$

In this formulation, we use $(\mathbf{r}_j, \mathbf{v}_j) \in \Sigma_j$ to mean that $\mathbf{v}_j$ is optimal for agent b_j given rewards $\mathbf{r}_j$:

$$\begin{aligned} (\mathbf{r}_j, \mathbf{v}_j) \in \Sigma_j &\iff r_{j,i} \in R, \quad \forall i \in \{1, \dots, n\}, \\ &\mathbf{v}_j = \operatorname{argmax} (\mathbf{u}_j + w_r \cdot \mathbf{r}_j)^T \cdot \mathbf{v}_j, \\ &\text{s.t. } \mathbf{c}_j^T \cdot \mathbf{v}_j \leq C_j. \end{aligned}$$

Our variant of the ADMM can be derived via the Augmented Lagrangian:

$$L_\rho = \sum_{j=1}^m \mathbf{s}^T \cdot \mathbf{v}_j + \lambda_j^T \cdot (\mathbf{r}_j - \mathbf{r}) + (\rho/2) \|\mathbf{r}_j - \mathbf{r}\|_2^2.$$

in which λ_j 's are Lagrangian multipliers, $\rho > 0$ is the penalty parameter. Our variant starts with an initial $\mathbf{r}_j^0, \mathbf{v}_j^0, \lambda_j^0$ and $\mathbf{r}^0$, and updates the Lagrangian in an alternating manner for T steps. At the k -th step, $(\mathbf{v}_j^{k+1}, \mathbf{r}_j^{k+1})$ and $\mathbf{r}^{k+1}$ are obtained by minimizing $L_\rho(\cdot)$ w.r.t. $(\mathbf{v}_j, \mathbf{r}_j)$ and $\mathbf{r}$, respectively. λ_j^{k+1} is updated by taking a subgradient step in the dual. The updates of OptIKA-L2T-ADMM are:

$$\begin{aligned} (\mathbf{v}_j^{k+1}, \mathbf{r}_j^{k+1}) = & \operatorname{argmin}_{(\mathbf{v}_j, \mathbf{r}_j) \in \Sigma_j} \mathbf{s}^T \mathbf{v}_j + \lambda_j^{kT} (\mathbf{r}_j - \mathbf{r}^k) \\ & + (\rho/2) \|\mathbf{r}_j - \mathbf{r}^k\|_2^2, \end{aligned} \quad (11)$$

$$\mathbf{r}^{k+1} = \frac{1}{m} \sum_{j=1}^m (1/\rho) \lambda_j^k + \mathbf{r}_j^{k+1}, \quad (12)$$

$$\lambda_j^{k+1} = \lambda_j^k + \rho(\mathbf{r}_j^{k+1} - \mathbf{r}^{k+1}). \quad (13)$$

The difference of our variant with classical ADMM is that we impose extra constraints $(\mathbf{v}_j, \mathbf{r}_j) \in \Sigma_j$ in the first optimization step in Equation 11. This makes it computationally hard. In practice, we solve it via MIP, using the three encodings as described above.³ However, the benefit of this algorithm is that the optimization problem for agent b_j is *localized*: it only involves variables and constraints for agent b_j herself, which represents a significant improvement over the previous algorithms, in which we need to consider all m agents all together in one encoding.

ADMM allows us to derive a series of interesting properties about the obtained solution. The Lagrange dual function $g(\{\lambda_j\})$ is defined as:

$$g(\{\lambda_j\}) = \inf_{\mathbf{r}, \mathbf{v}_j, \mathbf{r}_j: (\mathbf{v}_j, \mathbf{r}_j) \in \Sigma_j} L_\rho(\{\mathbf{r}_j\}, \{\mathbf{v}_j\}, \{\lambda_j\}, \mathbf{r}).$$

We can view OptIKA-L2T-ADMM as an alternating direction descend algorithm trying to find the optimum of the optimization problem $\max_{\{\lambda_j\}} g(\{\lambda_j\})$. The proofs of the following theorems are left to the supplementary materials [40].

THEOREM 2. (Weak Duality) *The optimal objective value to the problem $\max_{\{\lambda_j\}} g(\{\lambda_j\})$ is a lower bound on the optimal value of the global problem (OptIKA).*

THEOREM 3. *$g(\{\lambda_j\})$ is concave for $\{\lambda_j | j = 1 \dots m\}$.*

THEOREM 4. *If OptIKA-L2T-ADMM converges, then $\mathbf{r}_1, \dots, \mathbf{r}_m$ and $\mathbf{r}$ all converge to the same vector.*

³In the case of OptIKA-L2T-DP (or **greedy**), Σ_j is then a relaxed constraint set, which only has constraints specified by the dynamic programming (or greedy) encoding. We use a big-M notation to handle the quadratic term $\|\mathbf{r}_j - \mathbf{r}^k\|_2^2$.

3.2 Identification Problem

In practice, parameters governing agents' preferences, such as $\mathbf{u}_j$, w_r , are unknown to us. The identification problem therefore is to learn these parameters by observing agents' reactions under different reward schemes. In our setting, we use road distance as cost $c_{j,i}$ (which is the main factor for accessibility) and we learn each agent's budget C_j from the historical mean. The variables left to estimate are the intrinsic utilities $u_{j,i}$ and the elasticity of external rewards w_r . We further assume that the intrinsic utility $u_{j,i}$ is parameterized by a set of features: $u_{j,i} = \mathbf{w}_u^T \cdot \mathbf{f}_{j,i}$, in which $\mathbf{f}_{j,i}$ includes both personal features related to agent b_j and environmental features related to location i . We assume agents are rational, therefore, their choices should always maximize the overall utility. In other words, suppose one agent chooses location set L_j , then $\sum_{i \in L_j} \mathbf{w}_u^T \cdot \mathbf{f}_{j,i} + w_r \cdot r_i \geq \sum_{i \in L'} \mathbf{w}_u^T \cdot \mathbf{f}_{j,i} + w_r \cdot r_i$, holds for any other set of locations L' , when the total distance to reach all locations in L' is within the budget. The identification problem then corresponds to finding $(\mathbf{w}_u, w_r)$ to satisfy all inequalities of this type. Because of the trivial solution $\mathbf{w}_u = \mathbf{0}$, $w_r = 0$, we aim to maximize the margin:

$$\begin{aligned} \text{Min } & \|\mathbf{w}_u\|^2 + w_r^2, \\ \text{s.t. } & \sum_{i \in L_j} \mathbf{w}_u^T \cdot \mathbf{f}_{j,i} + w_r \cdot r_i \geq \sum_{i \in L'} \mathbf{w}_u^T \cdot \mathbf{f}_{j,i} + w_r \cdot r_i + \\ & \Phi(L_j, L'), \quad \forall L' : \sum_{i \in L'} c_{j,i} \leq C_j. \end{aligned} \quad (14)$$

Here $\Phi(L_j, L')$ is a loss function, which applies different levels of penalties to location set L' , depending on how similar L' and L_j are. We choose $\Phi(L_j, L') = |L_j \setminus L'| + |L' \setminus L_j|$ in the experiment. In practice, not all constraints shown in Equation 14 can be satisfied. Therefore, we introduce linear slack variables, and the whole identification problem becomes:

$$\begin{aligned} \text{Min } & \|\mathbf{w}_u\|^2 + w_r^2 + C \sum_{j=1}^m \xi_j, \\ \text{s.t. } & \sum_{i \in L_j} \mathbf{w}_u^T \cdot \mathbf{f}_{j,i} + w_r \cdot r_i \geq \sum_{i \in L'} \mathbf{w}_u^T \cdot \mathbf{f}_{j,i} + w_r \cdot r_i + \\ & \Phi(L_j, L') - \xi_j, \quad \forall L' : \sum_{i \in L'} c_{j,i} \leq C_j. \end{aligned} \quad (15)$$

This is a novel application of structural SVM [39]. As another contribution, we developed a modified delayed constraint generation approach to solve the optimization problem as shown in Equation 15, which involves solving knapsack-type problems for both the prediction and the separation problem within the structural SVM.

4. EXPERIMENTS

4.1 Algorithm Performance

We first compare algorithms assuming different levels of rationalities for the organizer and agents, on synthetically generated benchmarks, in which the initial number of visits $X_{0,i}$ is drawn from a geometric distribution in order to introduce some spatial bias, and other variables are drawn from uniform distributions. All the experiments are run using IBM CPLEX 12.6, on machines with a 12-core Intel x5690 3.46GHz CPU, and 48GB of memory. We implement the distributed version of the ADMM-based algorithms with 12 cores, in which each agent problem is allocated to one core.

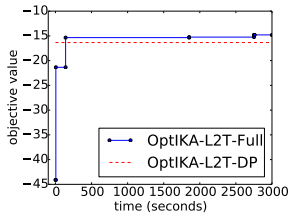


Figure 3: Comparison between OptIKA-L2T-Full and OptIKA-L2T-DP. (blue) Improvement of the objective function for OptIKA-L2T-Full over time. (red) Approximate solution value found by OptIKA-L2T-DP from solving a single MIP (very close to optimal, and much faster).

Method	Red.	δr	Method	δr
OptIKA-L2-Full	44%	0	OptIKA-L2-Full	0
OptIKA-L2-DP(50)	41%	1.36	OptIKA-L1-Full	1.20
OptIKA-L2-DP(100)	42%	1.13	OptIKA-L2-Full	0.74
OptIKA-L2-Greedy	41%	1.30		

Table 1: (Left) Comparison of different agents’ rationality level. *Red.* is the L2-norm reduction w.r.t. the non-incentive case, while δr is the average hamming distance of the reward vector w.r.t. the *Full*-rational case. (Right) Comparison of different organizer’s objectives, where δr is the average hamming dist. of the reward vector w.r.t. the L2 case.

Comparing Organizer’s Objectives & Agents’ Rationality Levels: We test our algorithms with 300 synthetic benchmarks. We fix the organizer’s goal (L2-norm), and consider the case where agents are planning with different levels of rationality. The left panel of Table 1 reports the reduction in terms of the organizer’s objective, and the mean hamming distance of the reward vectors obtained (i.e. the total number of locations in which the two reward vectors differ). Regarding the solution quality, the performance of the different approaches is similar in terms of the reduction in L2 and, while these approaches recommend 5 locations with positive rewards in the median case, the hamming distance between the reward vectors is barely more than 1. This suggests that the different models for the agents yield very similar results.

On the other hand, they largely differ in terms of computational complexity. When assuming full rationality of the agents, the row generation approach needs to solve multiple CPLEX instances iteratively. Fig. 3 depicts the running time for OptIKA-L2T-Full for one instance, compared with OptIKA-L2T-DP. As we can see, it takes the row generation algorithm a very long time to prove optimality, while OptIKA-L2T-DP finds a solution, close to optimal, and it is much faster. For a set of instances with 10 locations and up to 10 agents, the median completion time for the *Full* case is 1,251 seconds, while it corresponds to 80, 93 and 38 seconds for the single MIP in the DP case with $N_b = 50$, $N_b = 100$ and in the *Greedy* case, respectively.

Second, we study the impact of choosing different organizer’s objectives. As shown in the right panel of Table 1, the difference in terms of the solution quality is again very small. However, the running times vary significantly. The median completion times are 1,251, 11 and 10 seconds for L2, L1 and L2T objectives, respectively.

Convergence of ADMM: In order to measure how fast OptIKA-ADMM-L2T-DP converges, we first run OptIKA-ADMM-

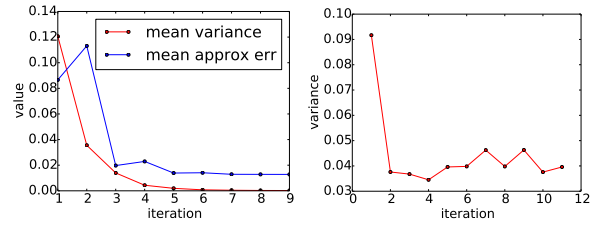


Figure 4: (Left) The mean variance and the relative error of OptIKA-ADMM-L2T-DP vs. iteration on small instances. (Right) Mean variance vs. iteration on a real *eBird* instance with 3,000 observers. ADMM converges very quickly.

Metric	SVM-struct	#Species	Popularity
Percentage Loss	3.9%	10.6%	8.2%
Utility Percentile	2.3%	46.0%	20.3%
Environmental Diff	2.0	5.9	4.9

Table 2: The Structural SVM model outperforms two other models on identifying people’s behavior (#Species: model based on estimated species num; Popularity: model based on location popularity).

L2T-DP for a set of small benchmarks with 20 or 30 agents and 20 locations. Although OptIKA-ADMM-L2T-DP works for problems of much larger scale, we still experiment with small benchmarks in order to compare it with non-decomposition based methods. In this experiment, the ADMM algorithm allocates one subproblem per agent. Because the main goal is to examine the decomposition method, each subproblem is solved by an OptIKA-L2T-DP module, and we compare the result with another OptIKA-L2T-DP which considers all agents at once. The two OptIKA-L2T-DP algorithms share a common discretization. The ρ for OptIKA-ADMM-L2T-DP is selected to be 1.

The blue line (top curve) in the left plot of Fig. 4 shows the relative error in the objective function as a function of the iteration number. The relative error is defined as $|S_{dp} - S_{admm}|/|S_{dp}|$, in which S_{dp} and S_{admm} are the objective values found by OptIKA-L2T-DP and OptIKA-ADMM-L2T-DP, respectively. The mean relative error is averaged among all benchmarks. As we can see, the error quickly drops from 10% to about 2% in only 3 iterations. At the same time, the red line shows how quickly the local copies $\mathbf{r}_j$ converge towards a common $\mathbf{r}$. For one benchmark, the variance is defined as: $\frac{1}{nm} \sum_{j=1}^m \|\bar{\mathbf{r}} - \mathbf{r}_j\|^2$, in which $\bar{\mathbf{r}}$ is the mean of $\mathbf{r}_1, \dots, \mathbf{r}_m$. The mean variance is taken among all benchmarks. As we can see, the variance drops to close to zero after 3 iterations.

Next we show the performance of OptIKA-ADMM-L2T-DP on an instance with 63 locations and 3,000 agents, which cannot be solved by non-decomposition methods at all. The agents’ behavior parameters come from real *eBird* data. We would like to emphasize that 3,000 is enough for real use, since there are in total 2,626 bird observers in New York State who submitted 3 or more observations in the past 10 years. The right plot of Fig. 4 shows the mean variance w.r.t. different iterations. Again we see that the local copies $\mathbf{r}_j$ almost converge to a common $\mathbf{r}$ in a few iterations.

4.2 Avicaching in eBird

eBird is a well-established citizen-science program for collecting bird observations. In its first years of existence, *eBird* mainly focused on appealing to birders to help address sci-

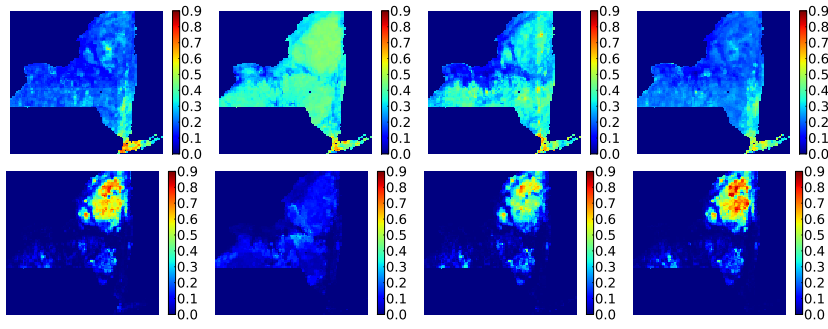


Figure 6: Heatmaps for the prediction of the White-throated Sparrow. (Upper 4 figures) Models for April. (Lower 4 figures) Models for July. A model trained on small, only 5% of the original data, but spatially uniform dataset (*Grid*, 2 in the leftmost column) has comparable accuracy with a model trained on the whole, big dataset that experts consider close to the ground truth (*Complete*, 2 in rightmost column), while other biased datasets have lower accuracy (*Urban*, 2nd column, *Random Subsample*, 3rd column).

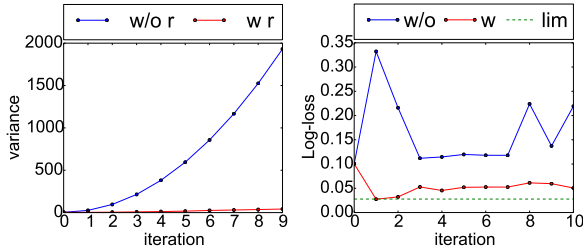


Figure 5: (Left) The change of spatial variance with and without incentives in the simulation study. (Right) The change of Log-loss to predict the occurrence of Horned Lark (with and without incentives). Dashed line is the performance limit.

Year	norm D_2	Treatment	norm D_2
2015	0.010	OptIKA	0.015
2014	0.016	B1: Inv-correlate #visits	0.021
2013	0.018	B2: Uniform-in-Avicache	0.017
		Manual (Expert's)	0.020

Table 3: (Left) Visits are more uniform (in normalized D_2) from April to August, 2015, when *Avicaching* is introduced, compared to previous years. (Right) Visits are more uniform under rewards designed by OptIKA against baseline B1 which assigns rewards inversely correlates to the number of visits to locations, B2 which assigns uniform rewards to all *Avicaching* locations, zero to others, and manually designed rewards (average over weeks of each treatment) in summer 2015.

ence objectives. The participation rates were disappointing. After 3 years, in order to make participation more fun and engaging, in the spirit of “friendly competition” and “cooperation”, *eBird* started providing tools to allow birders to track and rank their submissions (e.g., leaderboards by region, number of species, and number of checklists). This approach resulted in an exponential increase of submissions [35]. Nevertheless, like most citizen-science programs, *eBird* suffers from sampling bias. Birders tend to visit locations aligned with their preferences, leading to gaps in remote areas and areas perceived as uninteresting, as shown in Fig. 1.

In order to address this data bias, we gamified our methodology via a web-based application called *Avicaching*, explaining to birders that the goal of *Avicaching* is to “increase *eBird* data density in habitats that are generally under-represented by normal *eBirding*”. We deployed *Avicaching*

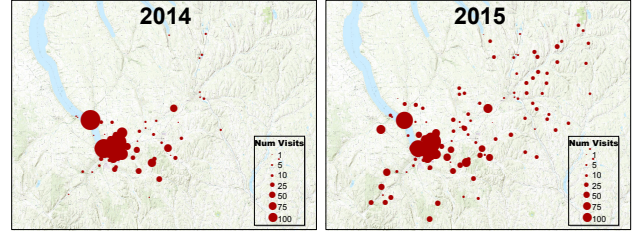


Figure 7: The number of *eBird* submissions in Tompkins and Cortland County in New York State. The circle sizes represent the number of submissions in each location. (Left) from Mar 28 to Oct 31, 2014 before *Avicaching*. (Right) from Mar 28 to Oct 31, 2015 when *Avicaching* is in the field. 19% effort is shifted to undersampled *Avicaching* locations.

as a pilot study in Tompkins and Cortland counties, NY, starting in March 2015. Tompkins is known for a high participation rate for *eBird*, while surprisingly, Cortland, a county adjacent to Tompkins, receives much fewer observations. We identified a set of locations with no prior observations and defined them as *Avicache* locations: bird watchers receive extra *avicache* points for every checklist submitted in those locations. The locations were selected around undercovered regions, emphasizing on important yet undersampled land types. We also ensure that all locations are publicly accessible. *Avicache* points have intrinsic values to bird watchers, because these points mark their contribution to science. In addition, participants have a chance to win a pair of binoculars from a lottery drawn based on their *avicache* points.

Pricing and Agents’ Model We update the *Avicaching* points every week. In the first few weeks, the allocation of points is manually assigned, based only on the number of previous visits to locations. This phase is designed to collect data to fit participants’ behavior model. After the initial phase, the points are assigned by the pricing algorithm (OptIKA-L1-DP). We fit the agents’ model using data from the two counties in 2015 (with *Avicaching* rewards), as well as data from the same season in 2013 and 2014 (without rewards). The results of the structural model are shown in Table 2, in which we predict the location set that people will visit per week. We randomly split all the data into a 90% training set and a 10% test set. The scores shown in the table are evaluated on the separate test set. The first measure is the percentage loss: $\frac{1}{n}(|L_{pred} \setminus L_{true}| + |L_{true} \setminus L_{pred}|)$, which is the difference between the predicted location set

with respect to the ground truth set. As shown in the table, **the mean percentage loss is merely 4%. This result is remarkably good, especially taking into account of the fact that we are modeling complex and noisy human behaviors.** We also look at how good our model is in terms of capturing people’s rationality. Ideally, we would like to see that our model always ranks the ground truth behavior the highest in terms of the utility score. Yet, this is impossible, because human beings occasionally take sub-optimal actions. In the utility percentile row, we show the percentile of the ground truth actions in terms of the utility scores among all valid actions. For example, the score 2.3% means that on average the utility scores of the ground truth actions are ranked at top 2.3% among all valid actions. Because the action set is big, we sample 10,000 location sets per test point. The low rank indicates that people are indeed motivated by the utilities defined in our model. Finally, the third row shows the difference of the environmental variables (NLCD values [19], normalized) between the predicted location set and the ground truth set. We compare our learned model with two other models. One chooses the set of locations which maximizes the estimated number of species (column #Species), and the other maximizes the total popularity of locations (column Popularity). These are the *two main* factors when planning a trip, according to expert opinions and birders’ surveys. In summary, our model is quite good at capturing agents’ preferences.

Field Results and Simulation We are delighted to see that people’s behaviors are changing with *Avicaching*.

1. Between Mar 28, 2015 and Aug 31, 2015, there have been 1,021 observations submitted from *Avicaching* locations, out of the 5,376 observations in total for these two counties: **19% birding effort has shifted from traditional locations to *Avicaching* locations, which received zero visits before. A few new birders also became more active, motivated by the *Avicaching* game.**
2. In terms of locations, Cortland, an undersampled county, received only 128 observations from April to August in 2013 and 2014 combined. This year during the same period of time, with *Avicaching*, it received **452 observations, over 3.5 times the total number of observations of the previous 2 years!**
3. Serious bird watchers are motivated to participate in *Avicaching*. **14 out of the Top 20 bird watchers in Tompkins and 15 out of the Top 20 bird watchers in Cortland (ranked by the number of species discovered since 2015) participated in *Avicaching*.** People who participated in *Avicaching* submitted 64% of total observations in Tompkins and Cortland, from April to August, 2015.
4. In terms of whether *Avicaching* is useful to motivate people to visit undersampled areas, we compare **OptIKA** against two baselines and a manually designed scheme. To eliminate time effects, we ensure that the results against baselines were all collected in summertime, with treatments interleaved. All baselines and **OptIKA** were given two weeks time. The numbers of locations receiving each level of rewards were kept the same for B1, **OptIKA**, and manual. The non-zero reward in B2 matches the mean of other treatments. Table 3 shows the comparison on the normalized D_2

score, which is $\frac{1}{n}||\mathbf{Y}-\bar{\mathbf{Y}}||_2/\bar{Y}$. The visits are more uniform in 2015, when *Avicaching* is introduced. Moreover, **OptIKA** wins against baselines in terms of uniformity. Figure 7 provides a visual confirmation on the map. The success of *Avicaching* is the combination of the gamification and the algorithm. It is difficult to isolate the algorithm’s contribution, because field implementation is time-consuming and we cannot afford to alienate the community with drastic or complicated experiments. The **OptIKA** algorithm is better in our experiment, but simpler algorithms may also work, especially at a small scale. However, they are likely to perform worse for new scenarios or over a large scale.

5. We further simulate, for a longer period, a set of virtual agents whose behaviors are learned from the *real bird watchers*. At the end of each round, we fit a predictive model based on all the data virtually collected so far, to see how much the species distribution model can be affected by agents’ shifts of exploration efforts. We use the collected data to predict the occurrence of the Horned Lark in Spring. The left plot of Fig. 5 illustrates the sample variance D_2 as a function of the number of iterations with and without extra incentives. The right plot of Fig. 5 shows the Log-loss of the fitted predictive model in the first few iterations. This simulation shows that ***under the Avicaching scheme, agents cover the area more uniformly*** than the no-incentive case, which leads to ***higher performance on a predictive model for bird occurrence***, with the same amount of effort devoted. More details of this simulation are in the supplementary materials [40].

Power of Uniform Sampling Finally, we also illustrate the benefit of incentivizing people to sample areas uniformly, by comparing the performance of a random forest classifier trained on four datasets, subsampled in different ways from the real *eBird* dataset. In Fig. 6, we show that the predictive model fit on a small, but spatially uniformly subsampled dataset is close to the ground truth, and outperforms the model fit on biased datasets. More details are in the supplementary materials [40].

5. CONCLUSION

We introduced a methodology to improve the scientific quality of data collected by citizen scientists, by providing incentives to shift their efforts to more crucial scientific tasks. We formulated the problem of Optimal Incentives for Knapsack Agents (**OptIKA**) as a two-stage game and provided novel algorithms on optimal reward design and on behavior modeling. We applied our methodology to *eBird* as a gamified application called *Avicaching*, deploying it in two NY counties. Our results show that our incentives are remarkably effective at steering the bird watchers’ efforts to explore under-sampled areas, which alleviates the data bias problem and improves species modeling.

Acknowledgments

We are thankful to Ronan Le Bras, the anonymous reviewers for comments, thousands of *eBird* participants, and the Cornell Lab of Ornithology for managing the database. This research was supported by National Science Foundation (0832782, 1522054, 1059284, 1356308), ARO grant W911-NF-14-1-0498, the Leon Levy Foundation and the Wolf Creek Foundation.

REFERENCES

- [1] G. Aggarwal, T. Feder, R. Motwani, and A. Zhu. Algorithms for multi-product pricing. In *ICALP*, 2004.
- [2] A. K. Agogino and K. Tumer. Analyzing and visualizing multiagent rewards in dynamic and stochastic environments. *Journal of Autonomous Agents and Multi-Agent Systems*, 17(2):320–338, 2008.
- [3] A. Anderson, D. P. Huttenlocher, J. M. Kleinberg, and J. Leskovec. Steering user behavior with badges. In *22nd Int'l World Wide Web Conference, WWW*, 2013.
- [4] D. F. Bacon, D. C. Parkes, Y. Chen, M. Rao, I. Kash, and M. Sridharan. Predicting your own effort. In *Proc. of the 11th Int'l Conf. on Autonomous Agents and Multiagent Systems-Volume 2*, pages 695–702, 2012.
- [5] R. Bonney, C. B. Cooper, J. Dickinson, S. Kelling, T. Phillips, K. V. Rosenberg, and J. Shirk. Citizen science: a developing tool for expanding science knowledge and scientific literacy. *BioScience*, 59(11):977–984, 2009.
- [6] S. P. Boyd, N. Parikh, E. Chu, B. Peleato, and J. Eckstein. Distributed optimization and statistical learning via the alternating direction method of multipliers. *Foundations and Trends in Machine Learning*, 3(1):1–122, 2011.
- [7] J. Bragg, Mausam, and D. S. Weld. Crowdsourcing multi-label classification for taxonomy creation. In *HCOMP*, 2013.
- [8] P. Briest, M. Hoefer, L. Gualà, and C. Ventre. On stackelberg pricing with computationally bounded consumers. In *WINE*, pages 42–54, 2009.
- [9] Y. Cai, C. Daskalakis, and C. H. Papadimitriou. Optimum statistical estimation with strategic data sources. *CoRR* 14, 2014.
- [10] X. Chen, Q. Lin, and D. Zhou. Optimistic knowledge gradient policy for optimal budget allocation in crowdsourcing. In *ICML*, 2013.
- [11] V. Chvatal. *Linear Programming*. W. H. Freeman and Company, New York, USA, 1983.
- [12] V. Conitzer and N. Garera. Learning algorithms for online principal-agent problems (and selling goods online). In *Proc. of the 23rd ICML*, 2006.
- [13] V. Conitzer and T. Sandholm. Computing the optimal strategy to commit to. In *Proc. 7th ACM Conference on Electronic Commerce (EC)*, pages 82–90, 2006.
- [14] U. Endriss, S. Kraus, J. Lang, and M. Wooldridge. Incentive engineering for boolean games. *IJCAI Proceedings-International Joint Conference on Artificial Intelligence*, 22(3):2602, 2011.
- [15] F. Fang, P. Stone, and M. Tambe. When security games go green: Designing defender strategies to prevent poaching and illegal fishing. In *IJCAI*, 2015.
- [16] P. Frazier, D. Kempe, J. M. Kleinberg, and R. Kleinberg. Incentivizing exploration. In *ACM Conference on Economics and Computation, EC*, pages 5–22, 2014.
- [17] V. Guruswami, J. D. Hartline, A. R. Karlin, D. Kempe, C. Kenyon, and F. McSherry. On profit-maximizing envy-free pricing. In *SODA*, pages 1164–1173, 2005.
- [18] J. D. Hartline and V. Koltun. Near-optimal pricing in near-linear time. In *Algorithms and Data Structures*, 9th Int'l Workshop, WADS, pages 422–431, 2005.
- [19] C. Homer, J. Dewitz, J. Fry, M. Coan, N. Hossain, C. Larson, N. Herold, A. Mckerrow, J. N. Vandriel, and J. Wickham. Completion of the 2001 National Land Cover Database for the conterminous United States. *Photogrammetric Engineering and Remote Sensing*, 73(4):337–341, 2007.
- [20] J. Hu and M. P. Wellman. Online learning about other agents in a dynamic multiagent system. In *Proceedings of the Second International Conference on Autonomous Agents*, AGENTS '98, 1998.
- [21] S. Jain, Y. Chen, and D. C. Parkes. Designing incentives for online question-and-answer forums. *Games and Economic Behavior*, 86:458–474, 2014.
- [22] R. Kawajiri, M. Shimosaka, and H. Kashima. Steered crowdsensing: Incentive design towards quality-oriented place-centric crowdsensing. In *UbiComp*, 2014.
- [23] H. Li, F. Tian, W. Chen, T. Qin, Z. Ma, and T. Liu. Generalization analysis for game-theoretic machine learning. In *AAAI*, 2015.
- [24] C. J. Lintott, K. Schawinski, A. Slosar, et al. Galaxy Zoo: morphologies derived from visual inspection of galaxies from the Sloan Digital Sky Survey. *Monthly Notices of the Royal Astronomical Society*, 389(3):1179–1189, Sept. 2008.
- [25] A. F. T. Martins, M. A. T. Figueiredo, P. M. Q. Aguiar, N. A. Smith, and E. P. Xing. An augmented lagrangian approach to constrained MAP inference. In *Proceedings of ICML*, pages 169–176, 2011.
- [26] P. Minder, S. Seuken, A. Bernstein, and M. Zollinger. Crowdmanager-combinatorial allocation and pricing of crowdsourcing tasks with time constraints. In *Workshop on Social Computing and User Generated Content in conjunction with ACM Conference on Electronic Commerce (ACM-EC 2012)*, 2012.
- [27] A. Y. Ng and S. J. Russell. Algorithms for inverse reinforcement learning. In *ICML*, pages 663–670, 2000.
- [28] P. Paruchuri, J. P. Pearce, J. Marecki, M. Tambe, F. Ordóñez, and S. Kraus. Playing games for security: an efficient exact algorithm for solving bayesian stackelberg games. In *AAMAS*, pages 895–902, 2008.
- [29] G. Radanovic and B. Faltings. Incentive schemes for participatory sensing. In *AAMAS*, 2015.
- [30] B. Settles. Active learning literature survey. *University of Wisconsin, Madison*, 52(55-66):11, 2010.
- [31] S. Shavell. Risk sharing and incentives in the principal and agent relationship. *The Bell Journal of Economics*, 1979.
- [32] Y. Singer and M. Mittal. Pricing mechanisms for crowdsourcing markets. In *Proceedings of the 22nd Internat. Conf. on World Wide Web (WWW)*, 2013.
- [33] A. Singla and A. Krause. Truthful incentives in crowdsourcing tasks using regret minimization mechanisms. In *Proceedings of the 22Nd International Conference on World Wide Web*, 2013.
- [34] A. Singla, M. Santoni, G. Bartók, P. Mukerji, M. Meenen, and A. Krause. Incentivizing users for balancing bike sharing systems. In *AAAI*, 2015.
- [35] B. L. Sullivan, J. L. Aycrigg, J. H. Barry, R. E. Bonney, N. Bruns, C. B. Cooper, T. Damoulas, A. A.

- Dhondt, T. Dietterich, A. Farnsworth, D. Fink, J. W. Fitzpatrick, T. Fredericks, J. Gerbracht, C. Gomes, W. M. Hochachka, M. J. Iliff, C. Lagoze, F. A. L. Sorte, M. Merrifield, W. Morris, T. B. Phillips, M. Reynolds, A. D. Rodewald, K. V. Rosenberg, N. M. Trautmann, A. Wiggins, D. W. Winkler, W.-K. Wong, C. L. Wood, J. Yu, and S. Kelling. The ebird enterprise: An integrated approach to development and application of citizen science. *Biological Conservation*, 169(0):31 – 40, 2014.
- [36] U. Syed, M. Bowling, and R. E. Schapire. Apprenticeship learning using linear programming. In *Proceedings of the 25th International Conference on Machine Learning*, ICML '08, 2008.
- [37] L. Tran-Thanh, T. D. Huynh, A. Rosenfeld, S. D. Ramchurn, and N. R. Jennings. Crowdsourcing complex workflows under budget constraints. In *Proceedings of the AAAI Conference*. AAAI, 2015.
- [38] L. Tran-Thanh, S. Stein, A. Rogers, and N. R. Jennings. Efficient crowdsourcing of unknown experts using bounded multi-armed bandits. *Artificial Intelligence*, 214:89–111, 2014.
- [39] I. Tsochantaridis, T. Joachims, T. Hofmann, and Y. Altun. Large margin methods for structured and interdependent output variables. *J. Mach. Learn. Res.*, 6:1453–1484, Dec. 2005.
- [40] Y. Xue, I. Davies, D. Fink, C. Wood, and C. P. Gomes. Avicaching: A two stage game for bias reduction in citizen science. *Technical Report*, 2016.
- [41] H. Zhang, E. Law, R. Miller, K. Gajos, D. Parkes, and E. Horvitz. Human computation tasks with global constraints. In *CHI*, 2012.
- [42] D. Zilli, O. Parson, G. V. Merrett, and A. Rogers. A hidden markov model-based acoustic cicada detector for crowdsourced smartphone biodiversity monitoring. *J. Artif. Intell. Res. (JAIR)*, 51:805–827, 2014.