

Data Coverage for Detecting Representation Bias in Image Datasets: A Crowdsourcing Approach

Melika Mousavi
mmousa7@uic.edu

University of Illinois Chicago

Nima Shahbazi
nshahb3@uic.edu

University of Illinois Chicago

Abolfazl Asudeh
asudeh@uic.edu

University of Illinois Chicago

ABSTRACT

Existing machine learning models have proven to fail when it comes to their performance for minority groups, mainly due to biases in data. In particular, datasets, especially social data, are often not representative of minorities. In this paper, we consider the problem of representation bias identification on image datasets without explicit attribute values. Using the notion of data coverage, we develop multiple crowdsourcing approaches. Our core approach, at a high level, is a divide and conquer algorithm that applies a search space pruning strategy to efficiently identify if a dataset misses proper coverage for a given group. We provide a different theoretical analysis of our algorithm, including a tight upper bound on its performance which guarantees its near-optimality. Using this algorithm as the core, we propose multiple heuristics to reduce the coverage detection cost across different cases with multiple intersectional/non-intersectional groups. We demonstrate how the pre-trained predictors are not reliable and hence not sufficient for detecting representation bias in the data. Finally, we adjust our core algorithm to utilize existing models for predicting image group(s) to minimize the coverage identification cost. We conduct extensive experiments, including live experiments on Amazon Mechanical Turk to validate our problem and evaluate our algorithms' performance.

Artifact Availability:

The source code, data, and/or other artifacts have been made available at <https://github.com/melimou/ImageDataCvgCrwd>.

1 INTRODUCTION

Tracing back machine bias to its source, there have been major efforts to identify different types [21, 37, 44] and sources [13, 16, 54] of bias in data. *Representation bias* [48], in particular, happens when a dataset fails to represent some parts of the target population [52]. Lack of representation from certain minority groups in data has caused many instances of machine bias and algorithmic unfairness in data-driven algorithms. For example, Facebook's ad algorithm excluding women from seeing specific jobs [27], or commercial gender classification systems from Microsoft, IBM, and Face++ that performed up to 35% worse on dark skin women compared to light skin men [9]. Similarly, the blink-detection feature of Nikon Cameras misclassified Asian eyes as being closed [46] due to a lack of representation for this group. We shall also demonstrate similar results in our experiments in § 6.4.

Recognizing the potential harms of representation bias, data coverage [1, 2, 4, 5, 29, 35, 39, 53] has been introduced to ensure proper representation of minority groups in datasets used for decision making and building advanced data science tools. At

a high level, a dataset has proper coverage for a given group if it contains at least a certain amount of samples belonging to that group. Despite the extensive work on detecting lack of data coverage, *existing work is limited to the tabular data*. On the contrary, many of the well-known incidents of representation bias causing machine bias, including all aforementioned examples, are in non-tabular contexts, such as multi-media or textual data.

Admitting the wide range of multimedia databases, with attributes of interest being in different forms and cardinalities, as our first attempt in this project, we consider *image data and a small number of low-cardinality categorical attributes*. Our choices are motivated by image data's popularity in data science tasks and the reported unfairness issues in the image application domains. Our assumption of the attributes of interest follows the fact that sensitive attributes such as race and gender are low-cardinality and non-ordinal, where each value such as black or female represents a specific demographic group.

It is common that image datasets lack explicit values for attributes of interest (such as gender or race). An image dataset is often a collection of images from different domains with little to no information about their domain and which groups they belong to. As a result, even studying coverage over low-cardinality and categorical attributes of interests is challenging in these cases. There are multiple directions one can seek to overcome such challenges. Considering a small number of categorical attributes of interest (such as race and gender), one can use off-the-shelf automated techniques, such as classifiers, to first label tuples with their demographic information¹. Then, relying on the predicted groups, apply the coverage detection techniques to identify the lack of coverage in data. However, as observed in our experiments, this approach fails, mainly due to the following issues:

- (1) (*Machine Bias*): while the objective of identifying lack of coverage is to minimize machine bias, using (problematic) off-the-shelf models will transfer their biases into the labeled data, causing bias in the evaluation of the dataset. For example, consider a gender-detection classifier. Due to the inherent issues in how the classifier has been trained (and the data it used), it may perform differently across different minority groups. For instance, in our experiments (Table 2), we observed that the precision of a gender classifier from a well-known face recognition framework such as DeepFace [47] for females can get as low as 8% for a given image dataset.
- (2) (*Lack of distribution generalizability*): Existing tools are trained using data that may come from a different application domain, following a different distribution, and hence may not perform well on the dataset to be evaluated. Let us consider the example of a gender-detection classifier once again. Suppose the classifier has been trained using the standard portray images with a solid background. One cannot expect the classifier to perform well on randomly taken images [32]. Note that applying transfer learning techniques to retrain the model

This project was supported in part by NSF 2107290.

© 2024 Copyright held by the owner/author(s). Published in Proceedings of the 27th International Conference on Extending Database Technology (EDBT), 25th March-28th March, 2024, ISBN 978-3-89318-091-2 on OpenProceedings.org. Distribution of this paper is permitted under the terms of the Creative Commons license CC-by-nc-nd 4.0.

¹In presence of accurate predictive models, *our algorithms utilize them* to identify coverage with minimum cost to verify the correctness of their results (see § 5).

using the dataset to be evaluated is not helpful since we cannot expect a model to identify the representation biases of the dataset it is trained on.

Considering the above issues with the existing data-driven tools and techniques, a promising approach to consider is crowdsourcing: to efficiently use human workers to identify a lack of coverage issues. Crowdsourcing is particularly promising for image data, for the tasks such as image labeling, which while being challenging for the machine, are considered "easy" for human-being to conduct with minimal error². Using crowdsourcing for labeling the images with their attributes of interest can potentially add human bias into the process. Fortunately, accurate and reliable crowdsourcing that minimizes individual errors and biases has been studied well in the literature. Aggregating the responses of multiple crowd workers [6, 15, 28], and profiling the crowd [55, 57, 58] are some of the known techniques proposed for this purpose. A baseline solution then can be designed as a two-step process: first ask the crowd to provide the attribute values for all images in the dataset. Then apply off-the-shelf coverage identification techniques [4] to detect the uncovered groups. Cost-effectiveness, however, is a major requirement in crowdsourcing frameworks such as Amazon Mechanical Turk (MTurk) since there usually is a cost associated with each crowd task. The proposed baseline solution is ineffective in such frameworks because, depending on the size of datasets, it may require a significant number of tasks, meaning a considerable cost to study coverage in a given dataset. Hence, in this paper we study the problem of *identifying the lack of coverage in an image dataset with the minimum number of crowd tasks*. Existing research on bias detection in image data sets is limited to [26] that proposes a crowd-sourcing workflow to facilitate discovering attributes with potential sampling bias (e.g., finding out the airplanes in an image data set are facing right). The high level idea in this paper is to show random samples of images from an input dataset to the crowd to identify common similarities, and then ask the crowd judge to verify the discovered statements (see § 7 for more details). In contrast, our objective in this paper is to detect representation bias in terms of the data coverage with respect to the given attributes of interest such as gender, instead of discovering attributes that may reflect potential sampling bias.

Summary of Contributions. We study coverage identification in image datasets using crowdsourcing. To the best of our knowledge, our paper is *the first to study data coverage in image datasets*. In summary, our contributions are the following:

- We propose a divide-and-conquer (d&c) algorithm to identify the coverage of a demographic group across an image dataset. To enable the development of our algorithm, we employ set-based crowd tasks, which have been utilized in various crowdsourcing studies [24, 33, 34, 36]. At a high level, our algorithm falls in the general class of group testing approaches, wherein the process of identifying certain objects is divided into tests on groups of items [17, 18]. While, following the same logic of the group testing approaches, the design details of our algorithm are problem specific, making its performance close to the lower bound on the maximum number of tasks. We prove this by a tight upper bound on the maximum number of tasks the algorithm generates.

- Using our d&c algorithm as the core, we propose efficient algorithms for coverage identification over different scenarios with multiple non-intersectional and intersectional groups. We further introduce practical heuristics for coverage identification by carefully aggregating the minority groups into the so-called "super-groups".
- In presence of pre-trained classification models that predict the group(s) an image belongs to, we adjust our core algorithm to utilize their prediction and minimize the coverage identification cost. In cases where the models accurately predict the group labels, our algorithm only generates a small number of tasks to verify the correctness of the results.
- We evaluate our algorithms using extensive experiments on real and synthetic settings. We run *live experiments on MTurk with real workers* to validate our proposal. Besides, our performance evaluation results verify our theoretical findings, confirming the effectiveness of our algorithms.

2 PRELIMINARIES

2.1 Data Model

We consider a dataset in form of a collection $\mathcal{D}$ of N objects, each being an image. For example, $\mathcal{D}$ can be a set of $N = 10,000$ human face images. We use the notation t_i to refer to the i -th object in $\mathcal{D}$. We consider the *no explicit attribute-value* model for the data. That is, $\mathcal{D}$ only contains the objects, while the objects are not annotated.

We assume objects are associated with at least one attribute of interest considered for identifying representation bias. Formally, we use $\mathbf{x} = \{x_1, \dots, x_d\}$ to specify the attributes of interest. Each attribute of interest is a categorical sensitive attribute such as race, gender, and age-group. Each attribute has a cardinality of two or more, specifying different non-overlapping (demographic) groups. For example, a binary attribute gender with values {male, female} partitions the individuals into two non-overlapping groups.

In particular, we consider three different scenarios with attribute and group models: The most simple scenario is the *single binary* sensitive attribute case where objects are associated with only one binary sensitive attribute. Many of the problematic representation bias cases that have been reported fall under this category. Examples of this type of attribute include skin-tone aggregated into a binary feature of fair and dark skin-tone [10]. Studies showed that the pulse oximeter devices have a questionable accuracy in measuring arterial oxyhemoglobin saturation in individuals with dark skin-tone [20], which proves it is imperative that skin-tone feature be taken into account in developing this type of device.

The immediate generalization is the *multiple non-intersectional* groups case where each object is associated with one sensitive attribute with cardinality larger than two. A non-binary attribute race with values {White, Black, Hispanic, Asian, Others}, or a multi-valued attribute gender with values {male, female, non-binary} are examples of this case.

The next level of generalization is the intersection of multiple attributes where each object is associated with more than one sensitive attribute that can be either binary or non-binary. An example of this case can be the intersection of race and gender, where each individual can be associated with one value from each of these attributes such as Asian female or Hispanic male.

²We do not make the assumption that human beings are completely reliable but rather suggest our framework for cases where human labels are more reliable.

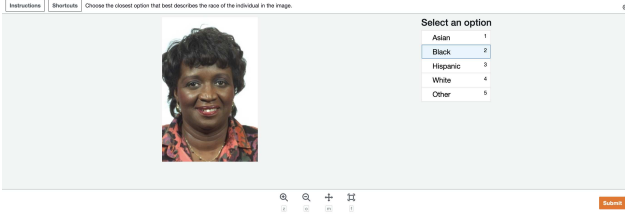


Figure 1: An example of point query to label a race attribute

2.2 Data Coverage

We use the notion of *data coverage* [4] to identify representation bias in a dataset $\mathcal{D}$. In particular, consider a dataset $\mathcal{D}$ with d attributes of interest $\mathbf{x}$, a count threshold τ (e.g. $\tau = 50$), and a subgroup $\mathbf{g}$ (e.g. $\{\text{gender}=\text{male}, \text{race}=\text{white}\}$) defined over $\mathbf{x}$. The dataset satisfies coverage over $\mathbf{g}$, if there are at least τ objects in $\mathcal{D}$, matching the subgroup $\mathbf{g}$ (e.g. there are more than 50 objects with $\text{gender}=\text{male}$ AND $\text{race}=\text{white}$).

For datasets with more than one attribute of interest ($d > 1$), *patterns* are used to specify the subgroups. A pattern P is a string of d values, where $P[i]$ is either a value from the domain of x_i , or it is “unspecified”, specified with X . For example, consider a dataset with three binary attributes of interest $\mathbf{x} = \{x_1, x_2, x_3\}$. The pattern $P = X01$ specifies all the tuples for which $x_2 = 0$ and $x_3 = 1$ (x_1 can have any value). Consider the universe of all patterns over a set of attributes $\mathbf{x}$. We say a pattern P is a parent of another pattern P' if (a) there exists exactly one attribute x_i on which P and P' are different, while (b) $P[i] = X$ (unspecified). Note that P in this case is a more general subgroup than P' , since all the objects matching P' also match P , but the vice-versa is not valid. A pattern P is a *maximal uncovered pattern* (MUP) if there are less than τ objects in $\mathcal{D}$ matching it, while all of its parents are covered. The lack of coverage in a dataset is identified by discovering all of its MUPs.

2.3 Crowdsourcing Model

A major challenge in studying coverage over multimedia data is that the objects are not annotated and we do not know the values on $\mathbf{x}$ beforehand. We use crowdsourcing to overcome this challenge. In crowdsourcing platforms such as MTurk or Crowdfunder, a *microtask* or a *Human Intelligence Task (HIT)* is a simple task that usually requires no domain-specific knowledge or expertise, has a clear description and a price which workers can accept and complete and get paid given their result is approved by the requester of the task.

Quality control and aggregation model. Quality of answers is a well-studied crowdsourcing challenge. One popular approach is to employ a redundancy-based strategy in which a single HIT is assigned to multiple workers and the correct answer (the *truth*) is inferred by aggregating the multiple answers. There are several studies on truth inference methods. The proposed techniques in this paper are *agnostic* to the choice of the crowdsourcing framework, quality control, and HIT aggregation model. In our experiments, we adopt the popular majority vote strategy [61] to get the truth and Qualification and Rating [14] as individual assessments to ensure a higher quality of answers from the crowd.

Query/ HIT model. We consider two types of queries:

- (1) *Point queries*: A point query is a request to provide a piece of information (attribute value) about a single object. The

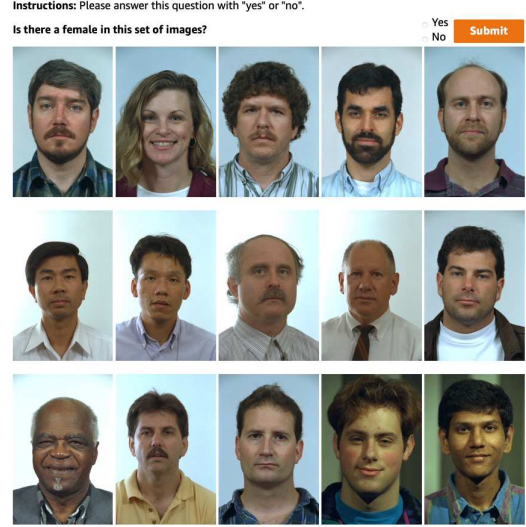


Figure 2: An example of set query about gender attribute

query itself can be either a yes-no question or providing one or more labels associated with the attributes of interest. In Figure 1 an example of a point query is demonstrated, where the worker is asked to provide the race of an individual.

- (2) *Set queries*: While a point query asks the crowd to provide specific attribute values for specific objects, the set queries are for the purpose of *verification*. A set query is a simple yes-no question about a given set of multimedia objects (image, video, etc.). That is to ask if the set contains at least one object belonging to a specific (sub-)group. For example, Figure 2 shows a set query asking the crowd to verify if the set contains any females. In practice, one may need to consider an upper bound on the number of objects in a set query for the query to be reasonable and the answers more accurate.

Pricing model. Pricing and incentive models for crowdsourcing frameworks have been extensively studied. Examples of such methods are *fixed price models*, *bidding models* [50], and *posted price models* [51]. Here, we adopt the fixed pricing model, that is all tasks have an equal cost. Therefore, our objective is to minimize the number of tasks required to finish the task of detecting coverage which in turn is in line with minimizing the total cost.

2.4 Problem Definition

Having explained the data model, data coverage, and the crowdsourcing model, we now define our problem as follows:

PROBLEM FORMULATION: Given an image dataset $\mathcal{D}$, the attributes of interest $\mathbf{x}$, and the coverage count threshold τ , identify the lack of coverage on $\mathcal{D}$ while minimizing the number of tasks required.

In particular, we study this problem under three settings: (1) single (minority) group, (2) multiple non-intersectional groups, and (3) intersectional groups, proposing efficient solutions tailored for each setting. At a high level, our main idea is to divide the dataset into subsets of a specified size, ask the crowd a question about the attributes of interest, and based on the crowd’s answer to the question, decide to divide that particular subset into two halves or prune it. Modeling each algorithm’s flow into a binary tree, we analyze the efficiency of our proposal, particularly when the dataset size is very large. We also study the

problem in the presence of predictive tools for labeling the data, utilizing them to effectively form the tasks, and minimizing the interaction with the crowd.

3 EFFICIENT COVERAGE IDENTIFICATION

We start our technical sections by designing a general algorithm that can be used for detecting coverage over different settings of sensitive attributes. In particular, *given a demographic (sub)group g , our goal in this section is to identify if g is uncovered.*

Before proposing the algorithm, however, let us consider the single binary sensitive attribute case to observe a challenge we need to address in our algorithm. In such a setting, the binary attribute partitions the data into two groups: the majority group (to which most of the objects belong) and the minority group. Let g_1 and g_2 be the majority (e.g. gender=male) and minority (e.g. gender=female) group, respectively.

Challenge. We observe that *verifying that g_1 is covered* can be done efficiently. To see why, suppose the coverage threshold is $\tau = 100$, i.e., a group is covered if there are at least 100 instances of it in the dataset. Assume the (majority) group g_1 contains $n_1 \gg 100$ objects in the dataset. In order to verify that g_1 is not uncovered, it is enough for the crowd to discover 100 of those objects, not the entire n_1 . Following this, $\Theta(\tau)$ provides a lower bound on the number of crowd tasks required for verifying the coverage for a given group. Still, this lower bound only holds for the groups that are covered, i.e. there at least τ of those in the dataset. Surprisingly, unlike the majority group, verifying that a minority group is indeed uncovered is cumbersome. This is because even though discovering τ objects from a group is enough for verifying that it is covered, one cannot verify a group is uncovered until there is a chance that the dataset might still have enough objects from that group. Thus, assuming a non-zero probability for each unlabeled object to belong to each group, *one might need to ask the crowd to label the entire dataset before one can confirm that a specific group is uncovered.*

3.1 Coverage Identification for a Given Group

Verifying that a minority group is uncovered is challenging. Our idea is to design a *divide and conquer algorithm* that, instead of point queries, uses *set queries* to iteratively eliminate subsets of data that *does not include any object from the given group*. At a high level, the algorithm asks a set query from the crowd, inquiring whether the selected set contains at least one object from the given group g (Figure 2). The user may provide two responses (yes/no). Interestingly, *in either case*, the response provides useful information that helps in the efficient study of the coverage:

- *The answer is “No”:* If the answer to a set query is no, it means the set does not include any object from the given group g . As a result, the algorithm can safely prune the set, asking no further questions about it. In particular, for a group that is not covered, one can expect to see no answers on large set queries helping to quickly prune a large portion of the dataset.
- *The answer is “Yes”:* A yes answer to a set query means that the set contains one or more objects from the group g . Therefore, the algorithm cannot prune the subset since it can have any number (larger than zero) of the objects in g . At the first glance, the queries with yes answers do not provide useful information as the algorithm cannot prune the subset (it needs to divide it to smaller subsets). However, a key observation is that, since the sets are disjoint, *the algorithm will require observing only a*

Algorithm 1 GROUP-COVERAGE

Input: Dataset $\mathcal{D}$, dataset size N , subset size upper bound n , coverage threshold τ , target group g

Output: Coverage of group g , the count lower-bound cnt

```

1:  $cnt \leftarrow 0$ ; Let  $Q$  = an empty queue
2: for  $i \leftarrow 0$  to  $N$  with step size  $n$  do // init roots of subtrees
3:    $root \leftarrow \text{node}(i, i + n)$ ;  $Q.add(root)$ 
4: while  $Q$  is not empty do
5:    $T \leftarrow Q.del\_top()$ 
6:    $(i, j) \leftarrow (T.b\_index, T.e\_index)$ 
7:    $ans \leftarrow \text{ASKQUESTION}(\{t_i, \dots, t_j\}, g)$  // set query
8:   if  $T.parent$  is null then
9:     if  $ans = \text{yes}$  then  $cnt \leftarrow cnt + 1$  else continue
10:  else
11:    if  $ans = \text{no}$  then // prune
12:      if  $T = T.parent.left$  then
13:         $T \leftarrow Q.del(T.parent.right)$  // observe from
Line 21 that left nodes are added first
14:      else continue
15:      if  $T.parent.checked$  then  $cnt \leftarrow cnt + 1$ 
16:      else  $T.parent.checked \leftarrow \text{true}$ 
17:      if  $cnt = \tau$  then return true // covered
18:      if  $j > i$  then //if  $setsize > 1$ 
19:         $T.left \leftarrow \text{node}(i, \lfloor \frac{i+j}{2} \rfloor)$ ;  $T.right \leftarrow \text{node}(\lfloor \frac{i+j}{2} \rfloor + 1, j)$ 
20:         $T.left.parent \leftarrow T$ ;  $T.right.parent \leftarrow T$ 
21:         $Q.add(T.left)$ ;  $Q.add(T.right)$ 
22: return false,  $cnt$  //uncovered

```

limited number of yes answers before it stops. That is because the number of such queries provides a *lower-bound* on the number of objects from g in the dataset. Hence, as soon as the lower-bound reaches to τ , the algorithm can stop, knowing that g is covered.

Based on the above observations, we design our divide and conquer algorithm as follows. We use a *binary tree* data structure to efficiently implement the algorithm. Each node in the binary tree has the following structure:

```

struct node:
  b_index    // the beginning index of the range
  e_index    // the end index of the range
  parent=null, // link to the parent node
  left=null,   // link to the left child
  right=null,  // link to the right child
  checked=false, // true if at least one of its
                // child nodes has returned a yes answer

```

Each node in the tree is associated with a set query containing the objects $\{t_{b_index}, \dots, t_{e_index}\} \in \mathcal{D}$. In addition, every node in the tree has pointers to its parent and children. Finally, every node contains a boolean variable “checked” (with the default value false) that is used for maintaining a lower-bound on the number of objects discovered from the target group g .

Using the tree data structure, Algorithm 1 shows the pseudo-code of our proposed algorithm for the single binary sensitive attribute case. The algorithm uses the variable cnt to maintain the lower-bound number of objects discovered from the target group g . Considering a maximum size of n for the set queries, the algorithm starts by partitioning the data into $\lfloor N/n \rfloor$ subsets, allocating each to a binary tree. Adding the roots of the trees to

the queues, the algorithm then iteratively removes a tree node from the queue until a lower-bound count of τ is achieved for the minority group g or it verifies that g is uncovered.

For every node T removed from the queue, the algorithm asks the crowd to verify if its corresponding set contains at least one instance belonging to g . Depending on the answer from the crowd, multiple situations can happen. If the node stands for the root of a binary tree and the response is yes, the algorithm has found at least one more object from g ; but if the response is no (Line 11), it can safely prune the entire set from the search space and continue with other sets. On the other hand, if T is not a root node, if the answer is no, it means the answer for the other child of its parent should be yes. That is because the parent node contained at least one object from g . Since T does not contain any such object, the other child of its parent should contain at least one. As a result, in Line 13, the algorithm replaces T with the other child of its parent, safely knowing the answer to a query to the new T is yes.

When the answer to a non-root node query is yes, the algorithm may or may not be able to increase the lower-bound cnt . Note that the algorithm has already associated at least one object from g to the parent of each non-root node (the nodes with no answers have been pruned). As a result, the lower bound on g gets increased if the answer to both children of a parent node is yes. We use the variable `checked` for this purpose. `checked` is true, if the answer to one of the children of a node is yes. Using this, when receiving a yes response, the algorithm increases the lower bound (Line 15) only if the `checked` variable of the parent of T is true. At any moment that the lower-bound reached the threshold τ , the algorithm stops marking g as covered. Finally, the algorithm breaks the yes nodes with set sizes larger than one in two halves, adding them to the queue. If after checking all nodes in the queue the threshold τ is not reached, g is uncovered.

Running Example. To better demonstrate GROUP-COVERAGE, let us consider a toy example with 16 images, where each image belongs to either group $\square$ or group $\triangle$. Suppose we would like to check if $\triangle$ is covered, while $\tau = 3$. Figure ?? shows the binary tree representation of the GROUP-COVERAGE algorithms' flow, while the root of the tree shows the entire image set. The answer to first query on the root is yes so the lower-bound value cnt gets updated to 1 and the images gets divided in two halves. The answer to both children of the root are also yes, so cnt gets updated to 2, while each set gets divided by half. The next level of the tree contains four set queries, each containing four images. Moving from left to right, the answer to the left-most query is no, therefore, (i) this set gets pruned and (ii) the algorithm *without issuing a new task* knows the answer to the second-from-left query is yes (otherwise its parent query could not be yes). The same situation happens for the two right nodes. Next, in the fourth level, the algorithm issues the first two queries from left and since the answer to both is yes, cnt gets updated to 3 and the algorithm stops since it reaches the coverage threshold τ . Note that in this example, the algorithm issues seven queries to the crowd before it stops.

3.2 Algorithm Analysis

LEMMA 3.1. (*Correctness*) The GROUP-COVERAGE algorithm successfully identifies if a group g is covered or not, i.e., if there are at least τ instances of g in the dataset $\mathcal{D}$.³

³Due to the space limitations, the proofs are provided in the technical report [40].

After correctness, let us now study the number of tasks GROUP-COVERAGE generate. Before studying the performance in general cases, let us consider two extreme cases, while assuming $N = n$, where (Case I) the answer to all questions is yes, and (Case II) there exists only one object in g .

Case I: Consider the cases where the answer to all set questions is yes, meaning that all set queries contain at least one object belonging to g . In this case, the set queries do not help prune the search space. As a result, the execution tree is a complete binary tree shown in Figure 3b. Note that the number of leaf nodes in this tree shows the value of cnt , knowing that the answers to those set queries are yes. On the other hand, the GROUP-COVERAGE algorithm stops when $cnt = \tau$. Therefore, the number of leaf nodes in the tree is at most τ . In a complete binary tree with τ leaves, there are $\tau - 1$ non-leaf nodes. As a result, the total number of tasks generated by the GROUP-COVERAGE algorithm, in this case, is $2\tau - 1 = \Theta(\tau)$.

Case II: Suppose there exists only one object from g in $\mathcal{D}$. In this case (Figure 3c), every level of the execution tree contains exactly one node while one of its children is no, while the other one is yes. Given that the root of the tree corresponds to a set of n objects, the leaf yes node (containing one object) is at depth $\log n$ in the tree. Hence, since every intermediate node has exactly two children, the number of tasks generated in this case is $\Theta(\log n)$.

THEOREM 3.2. Assuming the dataset size is $N = n$, or when there is no limit on the set query size, (i) the maximum number of tasks generated by the GROUP-COVERAGE algorithm is $\Theta(\tau \log n)$, (ii) the upper bound is tight.

LEMMA 3.3. (*Cost Analysis*) The maximum number of tasks generated by the GROUP-COVERAGE algorithm is $\Theta(\frac{N}{n} + \tau \log n)$.

Before concluding this section we would like to note that lower-bound on the maximum number of set queries an algorithm need to issue for deciding if g is covered is $\frac{N}{n}$. That simply is because $\frac{N}{n}$ queries are needed to include each object in $\mathcal{D}$ in at least one set query. In cases where g is uncovered, all objects should be queried to verify g is indeed uncovered. Comparing this lower-bound with the maximum number of tasks generated by the GROUP-COVERAGE algorithm, one can verify that GROUP-COVERAGE has only a small *additive* overhead of $\Theta(\tau \log n)$ from the lower-bound.

3.3 Coverage Identification using GROUP-COVERAGE

Recall that the algorithm presented so far considered the single binary attribute case. Particularly, given a group g , GROUP-COVERAGE efficiently interacts with the crowd to identify if g is covered. This algorithm can be applied for coverage identification for different scenarios of sensitive attributes, as followed:

3.3.1 Multiple Non-intersectional Groups. In a case where there is one attribute-of-interest x , let c be the cardinality of x . In this case, each value of x specifies a group g_i . To study coverage in this case, one needs to identify if each of the groups g_i is covered in $\mathcal{D}$. This can simply be done by running the GROUP-COVERAGE algorithm c times, where the i -th run checks if the group g_i is covered. As a result, following Lemma 3.3, the maximum number of tasks generated for this case is then $\Theta(c(\frac{N}{n} + \tau \log n))$.

3.3.2 Intersectional Groups. For cases with intersectional groups, the intersections of attribute values on $\mathbf{x} = \{x_1, \dots, x_d\}$ specify different subgroups, while the objective is to identify the uncovered region in form of *maximal uncovered patterns* (MUPs) [4].

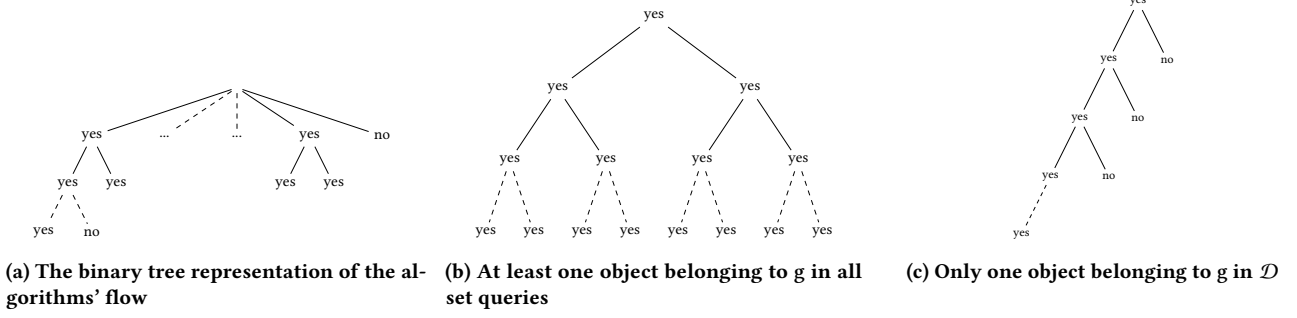


Figure 3: Various outcomes of the binary tree of the generated tasks

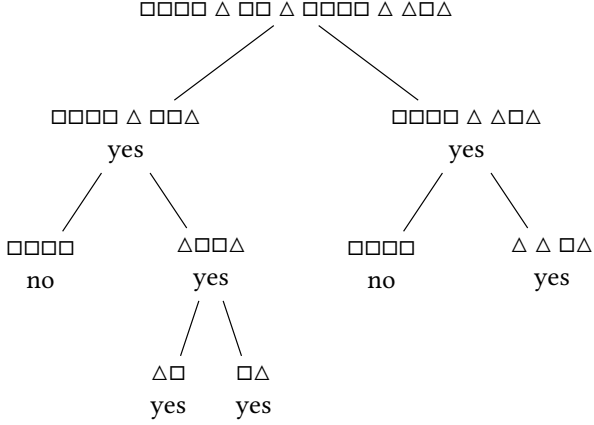


Figure 4: GROUP-COVERAGE: Running Example

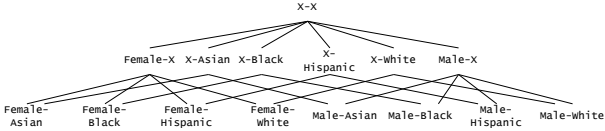


Figure 5: The pattern graph for x_1 =gender and x_2 =race

Consider the graph representation of patterns associated with the attributes x . A real-life example of this case is represented with gender and race attributes in Figure 5. The first level of the graph contains patterns, such as female- x or x -black with one attribute-value specified for them. A child of a node at level ℓ is a node at level $\ell + 1$ with one more attribute-value specified than its parent(s). For example, the level-2 pattern female-black (representing the subgroup of female black individuals) is a child of patterns female- x and x -black. An MUP is a pattern that is uncovered itself but all of its parents are covered. For example, female-black is an MUP if it is uncovered but both its parents female- x and x -black are covered.

In order to specify MUPs using the GROUP-COVERAGE algorithm, we consider the fully-specified subgroups (at maximum level), each being the intersection of d attribute values. The Cartesian product of the values of x_1 to x_d determine these subgroups. For example, in Figure 5 the nodes at level 2 of the graph (e.g. female-asian) show the fully-specified subgroups. Using $c_1, \dots, c_d$ as the cardinality of the attributes $x_1, \dots, x_d$, the number of fully-specified subgroups is $m = c_1 \times c_2 \times \dots \times c_d$. Let $g_1, \dots, g_m$ be the set of these subgroups. Finding the MUPs is then possible by combing the Pattern-combiner algorithm [4]

with the algorithm GROUP-COVERAGE. At a high level, the pattern-combiner starts from the bottom of the pattern graph (fully-specified subgroups), and counts the coverage for each of those, eliminating the covered nodes (with all of their parents). Then moving up in the graph, for the nodes that have not been pruned, it uses the counts of their children to check if those are covered, pruning the covered nodes and their parents while identifying the MUPs (uncovered nodes that all of their parents are pruned).

Algorithm 1 enables running pattern-combiner noting that GROUP-COVERAGE finds the *exact count* for an uncovered group. In such cases the lower-bound variable (*cnt* in Algorithm 1) contains the number of items belonging to the uncovered group. Additionally, pattern-combiner only needs the counts for the fully-specified subgroups that are uncovered as it already prunes the covered nodes.

4 PRACTICAL OPTIMIZATIONS BASED ON GROUP AGGREGATION

So far in previous sections, we designed an efficient algorithm to detect if a certain group is covered in $\mathcal{D}$. We further explained how this algorithm enables coverage identification for different cases of attributes of interest. In this section, we propose practical heuristics for cases with multiple non-intersectional and intersectional groups. In particular, we note that independently running the GROUP-COVERAGE algorithm for different groups specified by x to identify coverage, misses the opportunity to reuse the information collected during each run.

First, to avoid labeling objects multiple times we move the labeled objects from the unlabeled set $\mathcal{D}$ to the labeled set $\mathcal{L}$. Our main idea, however, is to form set queries that combine multiple demographic groups in one task – instead of one group. Specifically, consider a case where two or more demographic groups in the dataset are uncovered and there are very few items corresponding to them such that after combining them to one “super-group”, the result is still uncovered. For example, suppose the races Native American, Asian, and Middle Eastern are on the absolute minority that the summation of counts for all these three groups is less than the coverage threshold. In this case, instead of running the GROUP-COVERAGE algorithm once for each of them, we can run it for the super-group once. To do so, we change the set query to combine the individual groups with *OR* predicate.

Nevertheless, the challenge is that we do not have prior knowledge about the dataset $\mathcal{D}$ to form the super-groups. In order to obtain such information, we consider *estimating the counts using sampling*. To do so, we add a sampling phase at the beginning of our method, in which a small random subset of the dataset is presented to the crowd as *point queries*, with their task being to label the items. The results from this step gives us an estimation

Algorithm 2 MULTIPLE-COVERAGE

Input: Dataset $\mathcal{D}$, dataset size N , subset size upper bound n , coverage threshold τ , target groups $\mathbb{G}$, sample-size parameter $c = 2$

Output: Coverage of all groups in $\mathbb{G}$

```
1:  $\mathcal{D}, \mathcal{L} \leftarrow \text{LABELSAMPLES}(\mathcal{D}, \tau, c)$  //obtain  $c\tau$  random labels
2:  $\mathbb{G}_{agg} \leftarrow \text{AGGREGATE}(\mathcal{L}, \tau, \mathbb{G})$  //form the super-group
3:  $cov \leftarrow$  empty set
4: for  $\mathcal{G} \in \mathbb{G}_{agg}$  do:
5:    $\tau' \leftarrow \tau - \sum_{g \in \mathcal{G}} \mathcal{L}.\text{COUNT}(g)$ 
6:    $cvg, cnt \leftarrow \text{GROUP-COVERAGE}(\mathcal{D}, N, n, \tau', \mathcal{G})$ 
7:   if  $|\mathcal{G}| = 1$  then  $cov.\text{ADD}(\langle \mathcal{G}, cvg, cnt \rangle)$ ; continue
8:   if  $cvg = \text{true}$  then //if the super-group  $\mathcal{G}$  is covered
9:     for  $g \in \mathcal{G}$  do:
10:       $\tau' \leftarrow \tau - \mathcal{L}.\text{COUNT}(g)$ 
11:       $cvg, cnt \leftarrow \text{GROUP-COVERAGE}(\mathcal{D}, N, n, \tau', g)$ 
12:       $cov.\text{ADD}(\langle g, cvg, cnt \rangle)$ 
13:   else for  $g \in \mathcal{G}$  do:  $cov.\text{ADD}(\langle g, \text{false}, cnt \rangle)$ 
14: return  $cov$ 
```

of the demographic groups frequencies in the dataset. Based on this estimation, the algorithm will decide which groups to aggregate as super-groups. Next, we need to determine the sample-set size. Our intuition is that point queries are efficient for verifying the coverage of the majority group, since we expect to discover enough of those after $\Theta(\tau)$ point queries. Therefore, we can first issue the point queries to identify the majority group, while at the same time, we *piggyback* on the point query results to collect information about minorities and form the super-groups. Following this idea, we consider labeling a random subset of size $c\tau$ of $\mathcal{D}$ at the beginning of the algorithm, where c is a small constant (we found $c = 2$ as a good choice in our experiments). Note that in cases where initial point queries do not find at least τ objects from the majority group(s), the algorithm effectively identifies a subset of these instances and needs fewer queries to get to the coverage threshold. As a result, GROUP-COVERAGE rapidly stops detecting them as covered. One drawback of forming the super-groups is when the result for a set of super-groups is *covered*. In this case, we could not know whether one, two, or all groups are covered, and thus, we need to examine each separately. In other words, the aggregation strategy will incur a penalty cost when the super-groups are covered.

AGGREGATE function.⁴ (Line 2 of Algorithm 2) Let $\mathbb{G}$ be the list of groups in one attribute or the set of fully-specified subgroups in the intersection of multiple attributes. The count estimations based on the samples collected in the labeled set $\mathcal{L}$ are utilized to set up the super-groups for $\mathbb{G}$. We calculate the expected number of instances corresponding to each group in the dataset based on their occurrence in the sample. Let $\mathcal{L}.\text{COUNT}(g)$ return the number of objects in group g that belong to $\mathcal{L}$. Since point queries are selected randomly, $\mathcal{L}$ is a random sample of $\mathcal{D}$. Hence, the expected size of g is $\mathbb{E}[|g|] = N(\mathcal{L}.\text{COUNT}(g))/|\mathcal{L}|$. If the expected number is less than the coverage threshold, it is likely that this particular group is uncovered in the dataset and vice versa. Similarly, if the summation of the expected numbers for a set of groups is still less than τ , the super-group formed by merging them is expected to be uncovered. To use this idea for

forming the super-groups, we first sort the groups based on their count values in $\mathcal{L}$ ascending. This helps to put the minority groups nearby and merge them as super-groups. Then we make a pass over the sorted groups while maintaining the sum over their expected coverage. So far as the expected sums are less than τ , the algorithm keeps merging the groups into a super-group, and then it moves to the next super-group. It finally returns the list $\mathbb{G}_{agg}$ of the super-groups.

Multiple and Intersectional Groups Coverage. Using the idea of merging minority groups into super-groups, Algorithm 2 specifies the uncovered groups for the cases where there exist multiple, non-intersectional groups in a single attribute. In particular, for every group $\mathcal{G}$ in the set of aggregated groups $\mathbb{G}_{agg}$, the algorithm first specifies the number τ' of instances it needs to observe before it can conclude $\mathcal{G}$ is covered. Next it runs the GROUP-COVERAGE algorithm for identifying the coverage of $\mathcal{G}$ in $\mathcal{D}$. If $\mathcal{G}$ is not a super-group, the algorithm directly adds the coverage result of $\mathcal{G}$ to the output. For cases where $\mathcal{G}$ is a super-group, if $\mathcal{G}$ is covered, the algorithm fails to conclude if groups $g \in \mathcal{G}$ are covered or not. Therefore, it reruns the GROUP-COVERAGE algorithm for all of such individual groups g . On the other hand, for cases where the super-group $\mathcal{G}$ is uncovered, the algorithm concludes that all groups in $\mathcal{G}$ are uncovered.

We can take advantage of the above technique for multiple attributes, where we are interested in identifying the coverage of each of the individual and the intersectional groups. Figure 5 demonstrates the pattern graph for two attributes of gender and race. To solve the problem for this case, we take on a similar idea to the PATTERN-COMBINER algorithm[4]. The objective of the PATTERN-COMBINER algorithm is to find MUPs (maximal uncovered patterns) in a dataset. As mentioned before, an MUP is a pattern that is uncovered but all of its parents are covered. Consequently, all of the children of MUP are uncovered as well. For example, in the race and gender attributes case with $\tau = 50$, assuming that we find 15 instances of Asian-Female and 20 instances of Asian-Male, we can conclude that Asian group with total instances of 35 is uncovered as well. On the other hand, if there were 28 Asian-Female and 32 instances of Asian-Male, we could conclude that Asian group is covered, without any additional crowdsourced tasks.

We use this idea to reduce the problem of identifying the coverage of multiple attributes to identifying the coverage of the fully-specified subgroups at the maximum level. We can see that this problem can be easily transformed into solving it for multiple, non-intersectional groups. It is noteworthy that the aggregation process for this special case requires that only the nodes with the same parent be aggregated with each other. To this end, we used a flag (*multi*) in our aggregation algorithm to distinguish between the two cases. Having identified the coverage of the subgroups using the MULTIPLE-COVERAGE algorithm, we then proceed to identify the coverage of all other patterns in the upper levels. Algorithm 3 describes the details of the discussed method.

5 UTILIZING EXISTING PREDICTORS

While *solely* relying on ML models for coverage identification might be a problematic, in presence of accurate and well-developed models, we should be able to utilize them in order to reduce the coverage identification cost – i.e., the number of crowdsourcing tasks. In this section, we adjust our core algorithm for this purpose. In such settings, instead of calling GROUP-COVERAGE

⁴The pseudo codes of the functions are provided in the technical report [40].

Algorithm 3 INTERSECTIONAL-COVERAGE

Input: Dataset $\mathcal{D}$, dataset size N , subset size upper bound n , coverage threshold τ , set of attributes x

Output: Coverage for all individual and intersecting groups in x

```
1:  $\mathcal{L} \leftarrow \text{LABELSAMPLES}(\mathcal{D}, \tau)$ 
2: Let  $\mathbb{G}$  be the set of fully-specified sub-groups at the max level
3:  $\mathbb{G}_{agg} \leftarrow \text{AGGREGATE}(\mathcal{L}, \tau, \mathbb{G}, \text{multi} = \text{true})$ 
4:  $cov \leftarrow \text{MULTIPLE-COVERAGE}(\mathcal{D}, N, n, \tau, \mathbb{G}_{agg})$ 
5: Let  $Q$  = an empty queue
6: for  $\langle g, cvg, cnt \rangle \in cov$  do  $Q.ADD(\langle g, cvg, cnt \rangle)$ 
7: while  $Q$  is not empty do
8:    $T \leftarrow Q.del\_top()$ 
9:   if  $T.cvg = \text{true}$  then
10:     foreach  $p$  in  $T.ancestors$  do:  $Q.ADD(\langle p, \text{true}, T.cnt \rangle)$ 
11:   else
12:      $cnt \leftarrow 0$ 
13:     for  $\forall P \in T.parent.children - T$  do
14:        $cnt \leftarrow cnt + P.cnt$ ;  $Q.pop(P)$ 
15:     if  $cnt \geq \tau$  then  $Q.ADD(\langle T.parent, \text{true}, T.cnt \rangle)$ 
16:     else  $Q.ADD(\langle T.parent, \text{false}, T.cnt \rangle)$ 
```

in subsequent algorithms, one should call our classifier-aware algorithm (CLASSIFIER-COVERAGE – Algorithm 4).

Using pre-trained classifiers on the dataset gives us the prediction groups. The predicted labels can be used in order to reduce the prediction cost. However, *we still need to validate the correctness* of the results obtained by the classifier in order to determine the coverage of a given group. Suppose a gender classifier is applied to a dataset to return a set f as females. In order to identify the coverage for the female group, the main idea is to eliminate the falsely identified females (false positives), namely males, from f . To this end, we apply a similar idea to what we did in GROUP-COVERAGE algorithm: we create a crowdsourcing task with all the points in set f , and ask a *reverse question*: “Is there any individual in this set that is NOT female?”. If the answer is yes, it means there are some false positives in this set. Therefore, we take the divide-and-conquer approach by dividing the set into two halves and repeating the question until all false positive instances are eliminated.

A performance issue with this strategy, however, happens when the false positive rate of the classifier for the given group is high. In such settings the divide and conquer strategy keeps dividing the sets into fine granularity, resulting in many small set queries to ask. In such cases, labeling the data points in the female set using point queries to verify the classifier’s label might be more efficient. Following this observation, we propose a sampling phase to estimate the precision of the classifier on the positive group, i.e., the group on which we would like to verify coverage (females in this example). Similar to our proposed method in §4, we choose a small, random sample from the identified females set (in our experiments, we found that a sample size of 10% of the set classified as the given group would be a good choice). In the next step, we ask the crowd to label the sample using point queries and estimate the precision of the classifier, comparing the classifier label and the true label. We experimentally found that if 25% of the sample are false positives, it is safe to say that the precision of the classifier on the positive group is sub-optimal for the group coverage identification task. Based on the estimated precision of the classifier on the target group, our algorithm decides whether

Algorithm 4 CLASSIFIER-COVERAGE

Input: Dataset $\mathcal{D}$, dataset size N , subset size upper bound n , coverage threshold τ , target group g

Output: Coverage of g

```
1: Let  $G$  be the set of tuples in  $\mathcal{D}$  with predicted label  $g$ 
2: Let  $S$  be a sample of  $G$ 
3: for  $t \in S$  do  $\mathcal{L}.add(\langle t, \text{POINTQUERY}(t) \rangle)$ 
4: if  $\mathcal{L}.COUNT(g) > 0.25|G|$  then  $\mathcal{G} \leftarrow \text{PARTITION}(\mathcal{D}, G, n)$ 
5: else  $\mathcal{G} \leftarrow \text{LABEL}(\mathcal{D}, G, \tau)$ 
6: if  $|\mathcal{G}| \geq \tau$  then return true
7: else return GROUP-COVERAGE  $(\mathcal{D} - G, N, n, \tau - |G|, g)$ 
```

to eliminate the false positive objects using either partitioning or labeling strategies.

At the end of this process, we will have a set containing only objects associated with the queried group g (female in this example). If we already have at least τ instances of a group g , we can determine the coverage and stop the process. However, assuming that the number of discovered instances c' is less than τ , now we have to find at least $\tau - c'$, *false negative* instances of g , i.e., the instances that belong to g but are classified as not belonging to g before we can conclude g is covered (or to verify the number of false positive is less than $\tau - c'$ and hence g is uncovered). This can be done by applying the algorithm GROUP-COVERAGE on the set of objects classified as not g with the threshold $\tau - c'$. The details of the proposed method can be found in Algorithm 4.

Our experiments in § 6.3.2, utilizing various pre-trained classifiers show that our optimization for coverage identification for classified datasets can achieve remarkable performance.

6 EXPERIMENTS

In this section, we evaluate the performance of our methods for coverage identification of single and multiple demographic groups. Additionally, we explore our heuristic of coverage detection for a single demographic group (gender) on pre-processed image data with pre-trained classifiers. Finally, we deploy our system on a real crowdsourcing platform, MTurk, to explore the performance of the discussed methods with real workers.

6.1 Experiment Setup

The experiments were conducted using both synthetic and real datasets. The algorithms were implemented in Python.

- **Synthetic datasets:** To thoroughly assess our algorithms, we create synthetic data with a variety of distributions. As an example, in a single demographic group problem setting for gender, a data point can be either $\{ 'F', 'M' \}$. In these experiments, we simulate the behavior of the crowdworkers in answering queries.
- **Image datasets:** We use slices of FERET DB [43] and UTK-Face [60] image datasets for the purpose of experiments on MTurk and applying pre-trained classifiers on the data. FERET DB is a dataset of 14,126 images of 1199 individuals taken from a variety of angles. UTKFace is also an image dataset consisting of over 20,000 annotated facial images.

Evaluation Plan: We evaluate the performance of GROUP-COVERAGE and the optimizations for multiple non-intersectional and intersectional groups as well as classified data using pre-trained classifiers. We report the number of tasks required for each experiment setting. We use a straightforward baseline, called BASE-COVERAGE, to identify group coverage as a baseline method to compare our

algorithm’s results to. In this method, each task is created containing only one single data point and is asked about. Two outcomes are possible for this algorithm: either at some point, τ instances of objects associated with the group are identified and hence, the group is covered, or the algorithm goes through all the data and determines that the group is uncovered.

Default Values: To evaluate the performance of our algorithms, we fix the value of some of the introduced parameters in §3-5. The default value of c is 2, and we fix τ and n as 50 in all experiments except their respective parameter analysis experiments.

6.2 Summary of Results

Our proposed GROUP-COVERAGE algorithm achieved remarkable results in group coverage identification in our experiments. Even at the worst possible case (where the number of instances associated with the group in the dataset is close to the coverage threshold), both the synthetic and MTurk experiments with the real crowdworkers showed that our algorithm needs a significantly small number of tasks compared to the size of the dataset in order to achieve results. We also show that the upper bound discussed in § 3.2 is in fact tight. Our optimizations on multiple group cases (both intersectional and non-intersectional) proved to be effective in most cases compared to a brute force approach utilizing the GROUP-COVERAGE algorithm to identify the coverage of multiple groups. Additionally, the optimization of classified data using pre-trained classifiers achieved notable results in most cases, decreasing the number of required tasks by approximately 80% in some.

6.3 Proof of Concept

6.3.1 Amazon Mechanical Turk. In this experiment, we evaluate our proposed method and the crowd’s performance on a live crowdsourcing platform. We defined a female coverage identification task and published our HITs on Amazon Mechanical Turk. Each HIT contained a set of initial $n = 50$ images to present a reasonable workload to the crowd, with a maximum assignment to 3 workers for quality control. The layout of the HITs was designed as Figure 2. The crowd was asked to answer the questions with yes or no.

To examine how each method affects the final outcomes, we use three quality control techniques: Majority Vote, Qualification Test, and Rating [14]. We adopt the off-the-shelf *Majority Vote* as a group assessment to control the quality of outputs, in which we assigned the same HIT to 3 workers and took the majority vote as the truth. We also experiment with *Rating* and *Qualification Test* as two types of individual assessments to further verify that our workers have the appropriate skills for the tasks. For the rating assessment, we measure the workers’ performance using *NumberHITsApproved* and *PercentAssignmentsApproved* in MTurk and only allow workers who meet a certain criterion for approved HITs and assignments to perform our tasks. Additionally, we designed a qualification test to verify the workers’ competence before granting them access to our HITs as another method for quality control. This test has a similar layout to the original HITs, which also served the purpose of familiarizing the workers with the tasks. One interesting observation that was made is that despite the relatively low number of assignments per HIT, only 1.36% of the total 660 answers from the crowd in all experiments (in 220 HITs) were incorrect which did not affect the final result in each experiment run. Additionally, we did not detect a significant difference between the experiments with or without individual

FERET DB (females=215, males=1307)	GROUP-COVERAGE #HITs	BASE-COVERAGE #HITs	UPPER-BOUND #HITs
QC: Majority Vote	74	342	115
QC: Qualification Test, Majority Vote	75	386	115
QC: Rating (PercentAssignmentsApproved $\geq$ 95, NumberHITsApproved $\geq$ 100), Majority Vote	71	284	115

Table 1: Coverage identification for female on MTurk

assessments which further supports our idea about the tasks being relatively easy and straight-forward for the crowd and fault-proof to an extent.

We employed the fixed price model as our pricing strategy. In our first set of experiments, each HIT price was set to \$0.1. In the next experiments, we decreased the reward for each HIT to \$0.05. Interestingly, this did not discourage the workers to accept and complete our tasks. Overall, we paid a total of \$44.1 to the workers and \$8.82 to Amazon MTurk as service charges.

We used two different subsets of the FERET dataset. With $\tau = 50$, the results from each experiment setting can be found in Table 1 with a comparison to the baseline BASE-COVERAGE method and our theoretical upper bound ($\frac{N}{n} + \tau \log(n)$).

6.3.2 Existing predictors. To investigate the performance of the existing predictors on image datasets and our strategy to detect the coverage of a dataset utilizing these models, we ran a number of experiments using DeepFace [47] and another CNN-based facial demographic classifier [30] to predict the gender of the individual in datasets. We used a subset of images of unique individuals from the FERET dataset [43], and two 3,000-point subsets of UTKFace [60] with different distribution of females to evaluate the results for both covered and uncovered cases. We applied DeepFace with opencv and retinaface as the underlying face detectors. Next, we passed the predicted labels and the females set detected by the classifier as the input to our CLASSIFIER-COVERAGE algorithm. Our algorithm chooses between "partitioning" and "labeling" to eliminate false positives in the identified female set. It picks "partitioning" if the classifier is at least 25% precise on a sample of the set, and "labeling" otherwise. The results are reported in Table 2. We also include the coverage detection results using GROUP-COVERAGE to compare them with those of CLASSIFIER-COVERAGE.

For each experiment, the accuracy and the precision of the classifier on the dataset are calculated separately. As discussed before, the accuracy of some classifiers can vary by a large margin on different data. Moreover, a high level of accuracy does not necessarily guarantee reasonable precision in predicting the class of data points. More specifically, both of the classifiers had a relatively high accuracy in the classification task on the UTKFace dataset, but both also had a low precision in their prediction for the female demographic group, which further proves the fact that the performance of the existing predictor for sensitive groups is questionable in many cases.

In addition, the results show that our heuristics make the right decision for which false positive elimination strategy to use in most cases, leading to significantly fewer necessary tasks to get the result. Compared to our GROUP-COVERAGE algorithm used standalone, the proposed techniques can produce significantly better results in most cases and still competitive results in others.

dataset	Existing classifiers' performance		precision on female group	CLASSIFIER-COVERAGE		GROUP-COVERAGE #HITs
	classifier	accuracy		false positive elimination strategy	CLASSIFIER-COVERAGE #HITs	
FERET DB (females=403, males=591)	DeepFace (opencv)	79.57	99.5	Partition	14	80
	DeepFace (Retinaface)	84.1	100.0	Partition	17	80
	BaseCNN	64.48	59.19	Label	84	80
UTKFace (females=200, males=2800)	DeepFace (opencv)	93.56	52.02	Label	97	51
	DeepFace (retinaface)	94.16	56.15	Label	89	51
	BaseCNN	97.6	74.8	Label	69	51
UTKFace (females=20, males=2980)	DeepFace (opencv)	96.53	8.0	Label	134	221
	DeepFace (retinaface)	96.43	10.09	Label	143	221
	BaseCNN	97.6	21.59	Label	122	221

Table 2: The results of female group coverage detection on gender classified datasets

Before concluding this section, we would like to remind that our algorithms for coverage multiple non-intersectional and intersectional sensitive attributes run the GROUP-COVERAGE algorithm in their core. In other words, those algorithms can be viewed as issuing GROUP-COVERAGE multiple times. So we expect similar results for those cases, using the real crowd.

6.4 Downstream Tasks Consequences

In this experiment, we show how the lack of coverage may cause model performance disparity (unfairness) in the downstream tasks. In particular, using two computer vision tasks we observe that (a) lack of coverage may cause model performance disparity for an uncovered group, and (b) resolving the lack of coverage reduces the performance disparity for the uncovered groups.

6.4.1 Drowsiness Detection. Drowsiness detection systems are used to prevent accidents that are caused by drivers who fell asleep while driving. MRL eye dataset [22] is a large-scale human eye dataset containing infrared images captured in a variety of lighting conditions from 37 people. While some of the subjects in the dataset wear glasses, we intentionally disregarded the images of such subjects to make them uncovered and created a sample of size 26480 images belonging to two classes of *open* (14279 images) and *closed* eyes (12201 images). Following the same procedure, we generated 10 datasets and repeated each experiment 10 times, using different datasets. Using this as the training data, we built a CNN model, and evaluating the model, we observed that while it has an average overall accuracy of 91.5%, the average accuracy for the spectacled subjects is only 81%. Next, in order to confirm that the issue was due to the lack of coverage, we gradually added 20, 40, 60, 80, and 100 images from the uncovered region back to each of the classes of *open* and *closed* in the training data, and retrained and evaluated the model. The results are illustrated in Figure 6a. With an increase in the number of samples taken from the uncovered group, we observed a reduction in the accuracy/loss disparity of the model between a randomly sampled test set and a sample consisting exclusively of spectacled subjects.

6.4.2 Gender Detection. To further verify our proposal, we repeat a similar procedure, using UTKFace dataset. We extract a sample comprising 7055 face images from UTKFace such that each image belonged to a class of either *male* (3834 images) or *female* (3221 images). While extracting the sample, we intentionally picked the subjects only if they were Caucasian. Using this as the training data, we trained a CNN model to predict the gender of the subjects. We repeated this procedure on 10 different samples and observed that on average, there is a 1% disparity in the overall accuracy of the model versus the accuracy for the Black

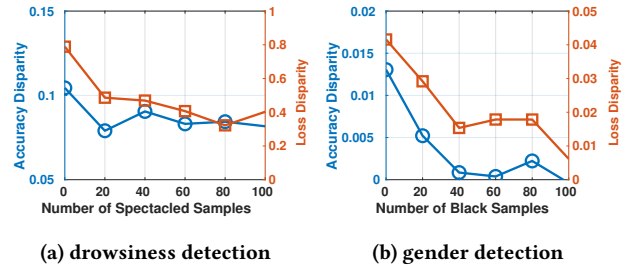


Figure 6: Effect of lack of coverage in the downstream tasks

subjects. Similar to the previous experiment, gradually increasing the number of Black subjects in the training data reduces the aforementioned disparity close to zero as seen in Figure 6b.

6.5 Performance Evaluation

In the following sections, we present the results of our experiments using the stated settings. First, we evaluate the performance of the GROUP-COVERAGE algorithm with varying parameters τ, n, N in §6.5.1. Next, we evaluate the optimizations for multiple non-intersectional and intersectional groups in §6.5.2.

6.5.1 GROUP-COVERAGE. To evaluate the performance of the GROUP-COVERAGE algorithm, we designed a simulation to reflect the procedure that the crowd would be presented with to carry out the tasks. The objective of the simulation is to first identify whether the dataset is covered with respect to a demographic group and determine the total number of tasks required to identify the coverage of a given group. For this purpose, assuming we are interested in identifying the coverage of female, we generate a dataset containing males and females and shuffle it randomly to prepare for the experiment. Each experiment with particular variables is run multiple times to better capture the effect of the dataset's underlying distribution on the results. In these sets of experiments, we study the impact of the scope of parameters on the end results.

Varying τ . First, we analyze the relationship between the coverage threshold and the number of females in the dataset and its impact on the number of necessary tasks to get the results. Figure 7a illustrates the number of required tasks when there exist $[0, 2\tau]$ items of the demographic group in the dataset. We have a dataset of size 100K, and we select the coverage threshold as 50. It can be observed that the largest number of queries is needed when the number of females (f) is close to τ . Conversely, the farther f gets from τ , the fewer tasks are required to get to a conclusion. This observation is consistent with the discussion in

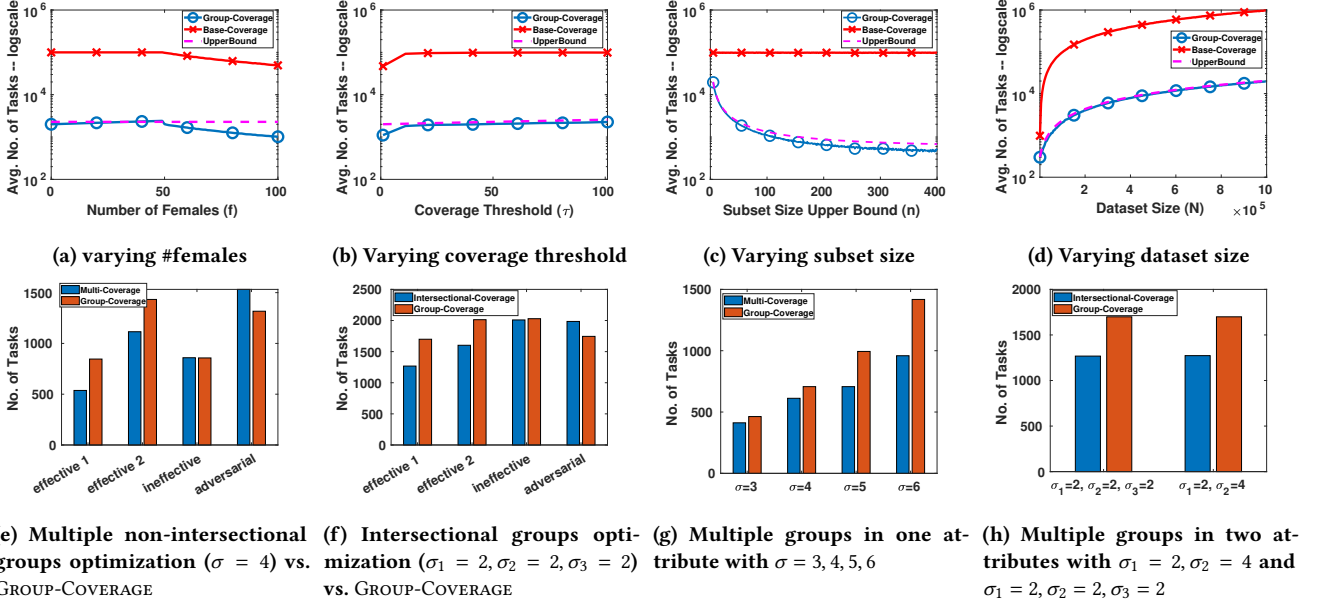


Figure 7: Performance evaluation for GROUP-COVERAGE, MULTIPLE-COVERAGE, and INTERSECTIONAL-COVERAGE algorithms

§3.2; with too few or too large quantities of f in the dataset, our algorithm’s results appear to be further from the upper bound.

Figure 7b shows the results of running the algorithm with different coverage thresholds. The coverage threshold is varied from 1 to 100 (0.001% to 0.1% of the dataset size) and there are exactly τ females at each run. Naturally, when the coverage threshold increases, the algorithm needs to cover more grounds to produce results. This also shows that the relationship between the coverage threshold and the cost is linear as discussed in §3.2. Note that the results in this figure demonstrate the case where $f = \tau$, which is the situation that requires the maximum number of tasks to get to a decision and is very close to the theoretical upper bound. **Varying n .** This experiment is designed to study the impact of the subsets size upper bound on the algorithm’s outcome. We set the coverage threshold to 50 in a dataset of size 100,000 while maintaining 50 instances of females in the dataset. In Figure 7c, we can see a substantial jump in the number of tasks when the subset size increases from around 10 to 20. Moreover, the result does not change significantly after that even with a notable increase in the subsets size. This confirms the logarithmic nature of the subsets size upper bound parameter in the algorithm.

When determining the set size, one must consider the crowd’s ability to identify the subject of the task at hand. While selecting larger n might lead to fewer required tasks, it is likely that we obtain less reliable answers from the crowd due to the large number of items presented all at once. Further, increasing the initial subset size will not significantly impact on the end results. **Varying N .** To assess our algorithm’s performance in a variety of datasets, we ran experiments for datasets of size 1K to 1M.

Figure 7d illustrates the results of the algorithm for varying dataset sizes. As expected, the number of required tasks to determine the results grows linearly with the size of the dataset, but never exceeds 6%. In other words, our results show that in practice, we can determine the group coverage for a dataset with tasks no more than 6% of the dataset size at the very worst case.

6.5.2 Optimizations for multiple groups. To evaluate the performance of our proposed method in identifying the coverage

for multiple groups, we take a similar approach as the previous section. We create a synthetic dataset comprising of data points that can correspond to $\sigma = 3, 4, 5, 6$ distinct demographic groups for the non-intersectional case, and two datasets, one with 2 attributes with cardinalities $\sigma_1 = 2, \sigma_2 = 4$ and the other with 3 binary attributes ($\sigma_1 = 2, \sigma_2 = 2, \sigma_3 = 2$) for the intersectional case. In a dataset of 10K points, with a threshold of 50 and a subset size of 50, we vary the number of items for each group in the dataset to simulate different combinations of settings and run the algorithm for each variation. Additionally, we run the GROUP-COVERAGE algorithm for each group independently to compare our results. In our experiments, we found out that while our heuristics on multiple groups can perform very well in some cases, it can also appear to be ineffective or worse than the brute force in some other. The results of the experiments for these cases are shown in Figures 7e and 7f.

Each of the bars defined as effective 1, effective 2, ... represent a different setting with respect to the number of instances associated with each group which is further described in Table 3.

The adversarial case, in which there are multiple uncovered groups with summation of items greater than τ in the dataset, it is likely that our heuristic fails in aggregating these groups into a super-group since the probability of having instances of these groups in the sample is significantly low. Thus, the super-group is covered and the algorithm needs to run for each of the subgroups individually. This imposes a penalty on the total number of tasks and makes it an adversarial case for our heuristic. To conclude, we can expect that our method works well or with little difference compared to brute force in some cases while failing in others.

Figure 7g shows the results for the MULTIPLE-COVERAGE algorithm for attributes with various cardinalities. Considering cases where our heuristic is effective, as the cardinality of the attribute increases, the total required tasks in MULTIPLE-COVERAGE grows more slowly than the brute force, resulting in a larger gap between the two methods as the cardinality increases.

Setting	Description
effective 1	3 uncovered minorities; their aggregated super-group is uncovered
effective 2	3 covered minorities
ineffective	2 uncovered and one covered minority
adversarial	3 uncovered minorities; their aggregated super-group is covered

Table 3: Experiment settings for multiple groups

Figure 7h represents the results of the INTERSECTIONAL-COVERAGE algorithm for two cases, one with 2 attributes with cardinalities $\sigma_1 = 2, \sigma_2 = 4$ and the other with 3 binary attributes ($\sigma_1 = 2, \sigma_2 = 2, \sigma_3 = 2$). As expected, with the same settings, the results for each of these cases are similar, with the number of fully-specified subgroups at the maximum level for both cases being equal. In other words, in the case of intersectional groups with multiple attributes, the only important feature is the cardinality of the attributes rather than the number of attributes.

7 RELATED WORKS

Crowd-sourcing for Bias Detection. [26] proposes a crowd-sourcing workflow to facilitate sampling bias discovery in visual datasets with the help of human-in-the-loop. This workflow takes a visual data set as an input and outputs a list of potential biases of the data set. There are three steps in this workflow. The first step is *Question Generation* and the crowd inspects random samples of images from the input dataset and describes their similarity using a question-answer pair. The next step is *Answer Collection* in which the crowd review separates random samples of images from the input dataset and provides answers to questions solicited from the previous step. Finally, in the third step called *Bias Judgement* the crowd judge whether statements of the visual dataset that are automatically generated using the questions and answers collected accurately reflect the real world.

Set queries. Set-based HITs, similar to our set queries, have been used in various crowdsourcing studies, including crowd powered data mining [33, 34]. [45] first introduced the idea of filtering a set of data based on a particular property using humans. Another example is [8], where for the purpose of top-k and group-by queries, the crowd is asked to answer *type* set question which has “yes” or “no” answer based on whether the data points in a set have the same type, which is similar to our notion of set queries on a target demographic group. Set queries have also been used in the crowd-sourced “count” operation: For instance, Marcus et al. [36] show a small batch of objects (images) to the crowd, asking them to estimate the number of items satisfying a specific constraint (e.g., photos with a car in them). Set queries are also popular in crowd-sourced clustering. For example, in [24] each worker views a small set of images as a HIT, where they are asked to provide a partial clustering of the set. Set queries have also been used for tasks such as crowd-sourced median finding [25], crowd-sourced planning [31], etc.

Group testing. Our approach in the GROUP-COVERAGE algorithm falls under the general category of group testing methods, where a task of identifying certain objects is broken up into tests on groups of items [18]. First proposed by Dorfman [17], group testing has been widely used across different domains [18],

with early applications such as detecting broken electrical circuits [11] with more recent applications in graphs [12], web databases [7], and even in Covid-19 detection [23]. Related work includes [19], which explores ways to perform efficient combinatorial group testing to identify up to d defective items from a set of n items using a reduced number of tests for practical set sizes. More generally, the class of divide and conquer (d&c) algorithms are popular in crowdsourcing. For example, [59] proposes a crowd-sourced d&c approach for sorting. Similarly, [42] proposes a crowdsourcing d&c approach for creating cross-lingual textual entailment corpora. Related work also includes crowd-sourcing d&c approaches for mobile platforms [3], paired comparisons [56], etc.

Coverage. The notion of data *coverage* has been studied across different settings [1, 2, 4, 5, 29, 35, 38, 48, 53]. With many angles to tackle, data coverage has been studied for datasets with discrete [4] and continuous [5] attributes populated in single or multiple [35] relations. Additionally, [1, 2, 49] (resp. [41]) use query rewriting (resp. data integration) to resolve representation bias. Existing works in data coverage have so far only focused on tabular data.

8 CONCLUSION

In this paper, we studied the problem of coverage identification in image data. This problem is motivated by the historical representation bias in various forms of data, and specifically the inefficiency of the existing supervised or unsupervised learning methods in performing equally well for minority groups on image data. We proposed an efficient algorithm to identify the coverage of a demographic group across the dataset and showed that the number of required tasks is optimal and close to the theoretical lower bound, and introduced practical heuristics to expand our solution for multiple non-intersectional or intersectional groups. We also presented an optimization method for detecting group coverage in datasets labeled by the existing predictors.

In this work, we focused on image data as a specific form of multimedia data. We hope to find equally efficient methods to identify data coverage in other forms of multimedia data such as video in our future work. In addition, our goal in this paper was mainly to minimize the cost of crowdsourcing by minimizing the total number of required tasks. We consider extending our techniques to support various pricing models as part of our future work.

9 RESEARCH ETHICS REVIEW STATEMENT

The research conducted in this study involving participants from MTurk ensured adherence to ethical principles and guidelines. Participants were provided with clear and comprehensive information such as the purpose, procedures, the type of the task, the compensation amount, and the expected time to complete the task. The participants were informed about the qualification screening. Informed consent was obtained from all participants before their engagement in the study. In addition to the platform anonymization of MTurk workers, we further ensured the privacy of participants by not collecting personal information.

REFERENCES

- [1] Chiara Accinelli, Barbara Catania, Giovanna Guerrini, and Simone Minisi. 2021. The impact of rewriting on coverage constraint satisfaction.. In *EDBT/ICDT Workshops*.
- [2] Chiara Accinelli, Simone Minisi, and Barbara Catania. 2020. Coverage-based Rewriting for Data Preparation.. In *EDBT/ICDT Workshops*.
- [3] Ariel Amato, Angel D Sappa, Alicia Fornés, Felipe Lumbleras, and Josep Lladós. 2013. Divide and conquer: atomizing and parallelizing a task in a mobile crowdsourcing platform. In *Proceedings of the 2nd ACM International Workshop on Crowdsourcing for Multimedia*. 21–22.
- [4] Abolfazl Asudeh, Zhongjun Jin, and HV Jagadish. 2019. Assessing and remediating coverage for a given dataset. In *2019 IEEE 35th International Conference on Data Engineering (ICDE)*. IEEE, 554–565.
- [5] Abolfazl Asudeh, Nima Shahbazi, Zhongjun Jin, and H. V. Jagadish. 2021. Identifying Insufficient Data Coverage for Ordinal Continuous-Valued Attributes. In *SIGMOD*. ACM.
- [6] Abolfazl Asudeh, Gensheng Zhang, Naeemul Hassan, Chengkai Li, and Gergely V. Zaruba. 2015. Crowdsourcing Pareto-Optimal Object Finding By Pairwise Comparisons. In *Proceedings of the 24th ACM International on Conference on Information and Knowledge Management (CIKM '15)*. Association for Computing Machinery, New York, NY, USA, 753–762. <https://doi.org/10.1145/2806416.2806451>
- [7] Abolfazl Asudeh, Nan Zhang, and Gautam Das. 2016. Query Reranking As A Service. *Proceedings of the VLDB Endowment* 9, 11 (2016).
- [8] Rubi Boim, Ohad Greenshpan, Tova Milo, Slava Novgorodov, Neoklis Polyzotis, and Wang-Chiew Tan. 2012. Asking the Right Questions in Crowd Data Sourcing. In *2012 IEEE 28th International Conference on Data Engineering*. 1261–1264. <https://doi.org/10.1109/ICDE.2012.122>
- [9] Joy Buolamwini and Timnit Gebru. 2018. Gender shades: Intersectional accuracy disparities in commercial gender classification. In *Conference on fairness, accountability and transparency*. PMLR, 77–91.
- [10] L. Elisa Celis and Vijay Keswani. 2020. Implicit diversity in image summarization. *Proceedings of the ACM on Human-Computer Interaction* 4, CSCW2 (2020), 1–28.
- [11] CC Chen and FK Hwang. 1989. Detecting and locating electrical shorts using group testing. *IEEE transactions on circuits and systems* 36, 8 (1989), 1113–1116.
- [12] Mahdi Cheraghchi, Amin Karbasi, Soheil Mohajer, and Venkatesh Saligrama. 2012. Graph-constrained group testing. *IEEE Transactions on Information Theory* 58, 1 (2012), 248–262.
- [13] Kate Crawford. 2013. The hidden biases in big data. *Harvard business review* 1, 4 (2013).
- [14] Florian Daniel, Pavel Kucherbaev, Cinzia Cappiello, Boualem Benatallah, and Mohammad Allahbakhsh. 2018. Quality control in crowdsourcing: A survey of quality attributes, assessment techniques, and assurance actions. *ACM Computing Surveys (CSUR)* 51, 1 (2018), 1–40.
- [15] Alexander Philip Dawid and Allan M Skene. 1979. Maximum likelihood estimation of observer error-rates using the EM algorithm. *Journal of the Royal Statistical Society: Series C (Applied Statistics)* 28, 1 (1979), 20–28.
- [16] Nicholas Diakopoulos. 2015. Algorithmic accountability: Journalistic investigation of computational power structures. *Digital journalism* 3, 3 (2015), 398–415.
- [17] Robert Dorfman. 1943. The detection of defective members of large populations. *The Annals of mathematical statistics* 14, 4 (1943), 436–440.
- [18] Dingzhu Du, Frank K Hwang, and Frank Hwang. 2000. *Combinatorial group testing and its applications*. Vol. 12. World Scientific.
- [19] David Eppstein, Michael T Goodrich, and Daniel S Hirschberg. 2007. Improved combinatorial group testing algorithms for real-world problem sizes. *SIAM J. Comput.* 36, 5 (2007), 1360–1375.
- [20] John R Feiner, John W Severinghaus, and Philip E Bickler. 2007. Dark skin decreases the accuracy of pulse oximeters at low oxygen saturation: the effects of oximeter probe type and gender. *Anesthesia & Analgesia* 105, 6 (2007), S18–S23.
- [21] Batya Friedman and Helen Nissenbaum. 1996. Bias in computer systems. *ACM Transactions on Information Systems (TOIS)* 14, 3 (1996), 330–347.
- [22] R. Fusek. 2018. Pupil localization using geodesic distance. *Lecture Notes in Computer Science (including subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)* 11241 LNCS (2018), 433–444. https://doi.org/10.1007/978-3-030-03801-4_38
- [23] Christian Gollier and Olivier Gossner. 2020. *Group testing against Covid-19*. Technical Report. EconPol Policy Brief.
- [24] Ryan Gomes, Peter Welinder, Andreas Krause, and Pietro Perona. 2011. Crowd-clustering. *Advances in neural information processing systems* 24 (2011).
- [25] Hannes Heikinheimo and Antti Ukkonen. 2013. The crowd-median algorithm. In *Proceedings of the AAAI Conference on Human Computation and Crowdsourcing*, Vol. 1. 69–77.
- [26] Xiao Hu, Haobo Wang, Anirudh Vegesana, Somesh Dube, Kaiwen Yu, Gore Kao, Shuo-Han Chen, Yung-Hsiang Lu, George K Thiruvathukal, and Ming Yin. 2020. Crowdsourcing Detection of Sampling Biases in Image Datasets. In *Proceedings of The Web Conference 2020*. 2955–2961.
- [27] Basileal Imana, Aleksandra Korolova, and John Heidemann. 2021. Auditing for discrimination in algorithms delivering job ads. In *Proceedings of the Web Conference 2021*. 3767–3778.
- [28] Panagiotis G Ipeirotis, Foster Provost, and Jing Wang. 2010. Quality management on amazon mechanical turk. In *Proceedings of the ACM SIGKDD workshop on human computation*. 64–67.
- [29] Zhongjun Jin, Mengjing Xu, Chenkai Sun, Abolfazl Asudeh, and HV Jagadish. 2020. MithraCoverage: A System for Investigating Population Bias for Intersectional Fairness. In *Proceedings of the 2020 ACM SIGMOD International Conference on Management of Data*. 2721–2724.
- [30] Shreyansh Joshi. 2020. Age and Gender Prediction from Facial Images. <https://github.com/ShreyanshJoshi/Facial-Demographics-using-CNN>.
- [31] Haim Kaplan, Ilia Lotosh, Tova Milo, and Slava Novgorodov. 2013. Answering planning queries with the crowd. *Proceedings of the VLDB Endowment* 6, 9 (2013), 697–708.
- [32] Bogdan Kulynych, Yao-Yuan Yang, Yaodong Yu, Jaroslaw Blasiok, and Preetum Nakkiran. 2022. What You See is What You Get: Distributional Generalization for Algorithm Design in Deep Learning. *arXiv preprint arXiv:2204.03230* (2022).
- [33] Guoliang Li, Jiannan Wang, Yudian Zheng, and Michael J Franklin. 2016. Crowdsourced data management: A survey. *IEEE Transactions on Knowledge and Data Engineering* 28, 9 (2016), 2296–2319.
- [34] Guoliang Li, Yudian Zheng, Ju Fan, Jiannan Wang, and Reynold Cheng. 2017. Crowdsourced data management: Overview and challenges. In *Proceedings of the 2017 ACM International Conference on Management of Data*. 1711–1716.
- [35] Yin Lin, Yifan Guan, Abolfazl Asudeh, and HV Jagadish. 2020. Identifying insufficient data coverage in databases with multiple relations. *Proceedings of the VLDB Endowment* 13, 12 (2020), 2229–2242.
- [36] Adam Marcus, David Karger, Samuel Madden, Robert Miller, and Sewoong Oh. 2012. Counting with the crowd. *Proceedings of the VLDB Endowment* 6, 2 (2012), 109–120.
- [37] Ninareh Mehrabi, Fred Morstatter, Nripsuta Saxena, Kristina Lerman, and Aram Galstyan. 2021. A survey on bias and fairness in machine learning. *ACM Computing Surveys (CSUR)* 54, 6 (2021), 1–35.
- [38] Yuval Moskovitch and HV Jagadish. 2020. Countata: dataset labeling using pattern counts. *Proceedings of the VLDB Endowment* 13, 12 (2020), 2829–2832.
- [39] Yuval Moskovitch and H. V. Jagadish. 2020. COUNTATA: Dataset Labeling Using Pattern Counts. *PVLDB* 13, 12 (2020), 2829–2832.
- [40] Melika Mousavi, Nima Shahbazi, and Abolfazl Asudeh. 2023. Data Coverage for Detecting Representation Bias in Image Datasets: A Crowdsourcing Approach. CoRR, abs/2306.13868, <https://arxiv.org/abs/2306.13868>.
- [41] Fatemeh Nargesian, Abolfazl Asudeh, and HV Jagadish. 2021. Tailoring data source distributions for fairness-aware data integration. *Proceedings of the VLDB Endowment* 14, 11 (2021), 2519–2532.
- [42] Matteo Negri, Luisa Bentivogli, Yashar Mehdad, Danilo Giampiccolo, and Alessandro Marchetti. 2011. Divide and conquer: Crowdsourcing the creation of cross-lingual textual entailment corpora. In *Proceedings of the 2011 Conference on Empirical Methods in Natural Language Processing*. 670–679.
- [43] National Institute of Standards and Technology. 2011. Face Recognition Technology Database (FERET DB). <https://www.nist.gov/itl/products-and-services/color-feret-database>.
- [44] Alexandra Olteanu, Carlos Castillo, Fernando Diaz, and Emre Kiciman. 2019. Social data: Biases, methodological pitfalls, and ethical boundaries. *Frontiers in Big Data* 2 (2019), 13.
- [45] Aditya G. Parameswaran, Hector Garcia-Molina, Hyunjung Park, Neoklis Polyzotis, Aditya Ramesh, and Jennifer Widom. 2012. CrowdScreen: Algorithms for Filtering Data with Humans. In *Proceedings of the 2012 ACM SIGMOD International Conference on Management of Data (SIGMOD '12)*. Association for Computing Machinery, New York, NY, USA, 361–372. <https://doi.org/10.1145/2213836.2213878>
- [46] Adam Rose. 2010. Are Face-Detection Cameras Racist? Time Business.
- [47] Sefik Ilkin Serengil and Alper Ozpinar. 2020. LightFace: A Hybrid Deep Face Recognition Framework. In *2020 Innovations in Intelligent Systems and Applications Conference (ASYU)*. IEEE, 23–27. <https://doi.org/10.1109/ASYU50717.2020.9259802>
- [48] Nima Shahbazi, Yin Lin, Abolfazl Asudeh, and HV Jagadish. 2023. Representation Bias in Data: A Survey on Identification and Resolution Techniques. *Comput. Surveys* (2023).
- [49] Suraj Shetiya, Ian P Swift, Abolfazl Asudeh, and Gautam Das. 2022. Fairness-aware range queries for selecting unbiased data. In *2022 IEEE 38th International Conference on Data Engineering (ICDE)*. IEEE, 1423–1436.
- [50] Yaron Singer and Manas Mittal. 2013. Pricing Mechanisms for Crowdsourcing Markets. In *Proceedings of the 22nd International Conference on World Wide Web (WWW '13)*. Association for Computing Machinery, New York, NY, USA, 1157–1166. <https://doi.org/10.1145/2488388.2488489>
- [51] Adish Singla and Andreas Krause. 2013. Truthful Incentives in Crowdsourcing Tasks Using Regret Minimization Mechanisms. In *Proceedings of the 22nd International Conference on World Wide Web (WWW '13)*. Association for Computing Machinery, New York, NY, USA, 1167–1178. <https://doi.org/10.1145/2488388.2488490>
- [52] Harini Suresh and John Gutttag. 2021. A framework for understanding sources of harm throughout the machine learning life cycle. In *Equity and Access in Algorithms, Mechanisms, and Optimization*. 1–9.
- [53] Ki Hyun Tae and Steven Euijong Whang. 2021. Slice tuner: A selective data acquisition framework for accurate and fair machine learning models. In *Proceedings of the 2021 International Conference on Management of Data*. 1771–1783.
- [54] Antonio Torralba and Alexei A Efros. 2011. Unbiased look at dataset bias. In *CVPR 2011*. IEEE, 1521–1528.

- [55] Jing Wang, Panagiotis G Ipeirotis, and Foster Provost. 2011. Managing crowdsourcing workers. In *The 2011 winter conference on business intelligence*. Citeseer, 10–12.
- [56] Ming-Hung Wang, Chia-Yuan Zhang, and Jia-Ru Song. 2023. CrowDC: A Divide-and-Conquer Approach for Paired Comparisons in Crowdsourcing. *arXiv preprint arXiv:2302.11722* (2023).
- [57] Peter Welinder, Steve Branson, Pietro Perona, and Serge Belongie. 2010. The multidimensional wisdom of crowds. *Advances in neural information processing systems* 23 (2010).
- [58] Jacob Whitehill, Ting-fan Wu, Jacob Bergsma, Javier Movellan, and Paul Ruvolo. 2009. Whose vote should count more: Optimal integration of labels from labelers of unknown expertise. *Advances in neural information processing systems* 22 (2009).
- [59] Haoqi Zhang, Eric Horvitz, Rob C Miller, and David C Parkes. 2011. Crowdsourcing general computation. In *ACM CHI 2011 Workshop on Crowdsourcing and Human Computation*. Citeseer, 1–5.
- [60] Zhifei Zhang, Yang Song, and Hairong Qi. 2017. Age Progression/Regression by Conditional Adversarial Autoencoder. In *IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*. IEEE.
- [61] Yudian Zheng, Guoliang Li, Yuanbing Li, Caihua Shan, and Reynold Cheng. 2017. Truth Inference in Crowdsourcing: Is the Problem Solved? *Proc. VLDB Endow.* 10, 5 (jan 2017), 541–552. <https://doi.org/10.14778/3055540.3055547>