



Detecting disparities in police deployments using dashcam data

Matt Franchi
mwf62@cornell.edu
Cornell Tech
New York, New York, USA

J.D. Zamfirescu-Pereira
University of California - Berkeley
Berkeley, USA
zamfi@berkeley.edu

Wendy Ju
Jacobs Technion-Cornell Institute, Cornell Tech
New York, USA

Emma Pierson
Jacobs Technion-Cornell Institute, Cornell Tech
New York, USA

ABSTRACT

Large-scale policing data is vital for detecting inequity in police behavior and policing algorithms. However, one important type of policing data remains largely unavailable within the United States: aggregated *police deployment data* capturing which neighborhoods have the heaviest police presences. Here we show that disparities in police deployment levels can be quantified by detecting police vehicles in dashcam images of public street scenes. Using a dataset of 24,803,854 dashcam images from rideshare drivers in New York City, we find that police vehicles can be detected with high accuracy (average precision 0.82, AUC 0.99) and identify 233,596 images which contain police vehicles. There is substantial inequality across neighborhoods in police vehicle deployment levels. The neighborhood with the highest deployment levels has almost 20 times higher levels than the neighborhood with the lowest. Two strikingly different types of areas experience high police vehicle deployments — 1) dense, higher-income, commercial areas and 2) lower-income neighborhoods with higher proportions of Black and Hispanic residents. We discuss the implications of these disparities for policing equity and for algorithms trained on policing data.

KEYWORDS

dashcam, object detection, policing

ACM Reference Format:

Matt Franchi, J.D. Zamfirescu-Pereira, Wendy Ju, and Emma Pierson. 2023. Detecting disparities in police deployments using dashcam data. In *2023 ACM Conference on Fairness, Accountability, and Transparency (FAccT '23)*, June 12–15, 2023, Chicago, IL, USA. ACM, New York, NY, USA, 11 pages. <https://doi.org/10.1145/3593013.3594020>

1 INTRODUCTION

How do citizens hold the police to account, ensuring that they are equitably protecting those they serve? A vital tool in the struggle for police accountability has been *data* shedding light on police

activity. Statistical analysis of police data has been used to document disparities in police stops, searches, use of force, and respect [2, 25, 27–29, 52, 53, 64, 73]. These statistical analyses have been instrumental to driving policy change and police reform: examples include a reduction in random police searches in Los Angeles following a statistical investigation documenting racial bias [9, 56] and a consent decree in Chicago following a Department of Justice investigation documenting numerous rights violations [71]. In recent years, policing datasets have become more widely available to the public [34, 53, 65], allowing activists, policymakers, and researchers to examine a range of policing behavior, including stops, searches, and arrests.

However, there is still one type of data that remains largely unavailable to the public within the United States: information on where the police are deployed in the first place. This data is crucial for several reasons. First, statistical analyses of police deployments have revealed inefficiencies or inequities [49, 50] in which some areas have disproportionate police deployments. Second, disparities in police deployments create biases in downstream outcomes like arrests. If an area has a heavier police presence, crimes are more likely to result in arrests. Thus, more heavily policed neighborhoods may appear to be more prone to crime, when they are simply more prone to *observed crime*; similarly, residents of these neighborhoods may appear to be more likely to commit a crime, when in fact they are simply more likely to be arrested for it. This bias has been observed in practice: for example, African Americans are far more likely to be *arrested* for marijuana use than are white Americans, even though surveys of drug use do not show the same racial disparities [57]. This bias could also propagate into *algorithms* used in policing and criminal justice — including pretrial risk assessments [15, 16, 42] and predictive policing algorithms [20, 45, 54, 58, 62] — since these algorithms can use arrests, or outcomes downstream from arrests like convictions, as input.

Deployment data are thus essential for monitoring disparities or inefficiencies in policing and detecting potential biases in algorithms. Even *aggregated* deployment data — grouped by neighborhood or demographic — would be useful towards this end. In spite of this, police deployment data remains largely unavailable to the public within the United States. Addressing this need, we introduce a novel methodology for monitoring police deployments: detecting police vehicles in dashcam images of public street scenes. Using a dataset of more than 24 million dashcam images collected throughout New York City in 2020, we annotate a training dataset of 9,449 images for presence of police vehicles, and train a deep

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for components of this work owned by others than the author(s) must be honored. Abstracting with credit is permitted. To copy otherwise, or republish, to post on servers or to redistribute to lists, requires prior specific permission and/or a fee. Request permissions from permissions@acm.org.

FAccT '23, June 12–15, 2023, Chicago, IL, USA

© 2023 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 979-8-4007-0192-4/23/06...\$15.00
<https://doi.org/10.1145/3593013.3594020>

learning model to identify police vehicles with high accuracy (average precision 0.82; AUC 0.99). We use this model to identify 233,596 dashcam images which contain police vehicles. We develop a framework for analyzing inequality in police deployment levels across neighborhoods — that is, the probability an image within a given neighborhood contains a police vehicle — which compensates for several data and model biases. Our analysis reveals substantial disparities in police deployments: the neighborhood with the highest police deployments has almost 20 times higher deployments than the neighborhood with lowest deployments. Two very different types of areas experience high police deployments: dense, high-income, commercial areas; and areas with higher proportions of Black and Hispanic residents and lower incomes. We discuss the implications of these disparities for policing and algorithmic equity. We also discuss the residual potential biases in the data which cannot be removed by our extensive debiasing pipeline, and argue that this fact suggests that the police should simply release more reliable and straightforward aggregated deployment data.

2 RELATED WORK

2.1 Auditing policing practices

An extensive academic literature uses large-scale policing data to statistically audit policing for efficiency and equity [1, 2, 25, 27–30, 45, 52, 53, 59, 64, 73, 75]. (Here and throughout the paper, we focus our discussion on policing within the United States because policing practices are highly heterogeneous across countries and our data is from New York City.) The existing academic literature largely analyzes outcomes *downstream* of police deployment patterns: that is, it studies not where police *go*, but what they *do once they get there*. For example, many papers [27, 28, 30, 52, 53, 59, 75] quantify disparities in how likely the police are to stop drivers of different races. Work has also studied post-stop outcomes: for example, [29] studies racial disparities in whether police grant leniency to drivers pulled over for speeding; [2, 36, 52, 53, 64] study racial disparities in police searches; and [73] studies racial disparities in how respectful officers are to drivers of different races. Work also quantifies police misconduct [31]; use of force [25]; police killings [21]; and downstream effects of police violence [1].

Outside of academia, there have also been extensive journalistic, activist, and legal efforts to quantify disparities in policing. Journalistic efforts include an investigation by the Los Angeles Times [10, 56] which documented racial bias in policing practices, ultimately resulting in a curtailment of random vehicle searches in an effort to reduce racial bias [9]. Closer to our own work, Mueller and Baker [49] analyzed confidential data on detective deployments within New York City and found racial and socioeconomic inequality in whether neighborhoods had enough homicide detectives to meet their needs. This report was followed by calls for increased transparency and more equitable deployments [50], testifying to the importance of police deployment data. Activists have also compiled large datasets on police activity [65]. Finally, many legal efforts against the police have leveraged large-scale policing data [67, 71].

Other efforts have quantified disparities in policing using datasets other than large-scale data from the police themselves. For example, journalists have analyzed footage of individual police encounters [66]; researchers have also interviewed individuals about their

experiences with police [22] and conducted analyses of survey data on experiences with police [19].

2.2 Effect of deployment disparities on algorithmic bias

Many algorithms used in policing and criminal justice use data from outcomes downstream from police deployments, including arrests and convictions. Examples include pretrial risk assessments [15, 16, 42] and predictive policing algorithms [20, 45, 54, 58, 62].¹ Thus, disparities in police deployments can bias the data these algorithms use [16, 20, 45, 54, 58, 62]. For example, Lum and Isaac [45] argue that disparities in police deployments could result in increasingly severe biases over time in predictive policing algorithms trained on arrest data. In their model, heavy police deployments in areas result in higher arrest levels at the same crime levels; consequently, the algorithm learns that areas with heavy police deployments are riskier than they really are, and recommends that even more police be deployed there; this further biases the arrest rates which are fed back into the algorithm, exacerbating inequality over time.

2.3 Analysis of public street scene datasets

Streetview imagery datasets – for example, Google Street View, Microsoft Bing Maps, Baidu Total View, and Tencent Street View – are increasingly used for urban analytics. Research has used such datasets to count trees [68, 77], utility poles [40], traffic signs [8, 69], bike racks [46] and manholes [72]. Streetview data has also been used for computational social science purposes: for example, estimating the demographic makeup of neighborhoods [26] or quantifying human political stereotyping [79].

More recently, *dashcam datasets* collected by Automotive Original Equipment Manufacturers (OEMs) and add-on manufacturers such as Nexar or Mobileye that provide cloud-connected dashcam products have enabled street view image datasets to be collected at a greater spatial scale and temporal density [78] (in contrast to Streetview data, which may be collected at much lower temporal density). This type of data has been used to assess pavement quality and classify road surface types [17], detect lanes or curbs [80], or detect potholes and other defects [60].

Dashcam data has a unique combination of properties that make it apt for the automated detection of objects in dense, feature-rich environments. Dashcam data is frame-by-frame, inexpensive to gather, not fixed in terms of camera perspective, and capable of depicting pedestrians, other vehicles, and small objects in detail. Dashcam datasets have been used to characterize when and where people were out in New York City during different phases of curfew and social distancing policy [12, 13, 35]. This work points to the possibility of using the density of the data from dashcam datasets to look at more dynamic aspects of urban environments—the number of people and cars, the density of traffic, and the prevalent activities [13].

Closest to our own work in this genre is Sheng et al. [63], which uses Streetview data to quantify how the density of surveillance cameras varies across and within cities. However, our work differs in important ways: most notably, we study police deployment, not

¹We note that the implementation details of many of these algorithms remain murky, and not all of them rely on arrest or conviction data.

surveillance cameras. Overall, the works are complementary: Sheng et al. [63] analyzes surveillance cameras across 16 cities at lower spatial and temporal resolution, and we analyze police deployments across a single city at high resolution.

3 DATASET

We estimate the total number and spatial distribution of police vehicles visible in public street scenes in New York City using a large dataset of dashcam images collected throughout 2020. In this section, we describe the dashcam dataset. In the next section, we describe our pipeline for analyzing it.

3.1 Dataset details

Our dataset is provided by Nexar, a company which provides ride-sharing (Uber, Lyft, NYC Taxi, etc) drivers with dashboard cameras (dashcams) to record their drives (which can be useful if, for example, an accident occurs). The dataset consists of 24,803,854 images taken throughout the five boroughs² of New York City between March 4 2020 and November 15 2020. Each image is 1280 x 720 pixels. We remove a small number of duplicate images (less than 0.01% of the overall dataset) with identical latitudes, longitudes, and timestamps.

3.2 Geographic and temporal coverage

Data was provided to us by Nexar in two phases. Phase 1 consists of 3,987,835 images sampled prior to September 1 2020, and is extremely geographically and temporally skewed. Geographically, it is concentrated within the boroughs of Manhattan and Brooklyn, and does not contain data from the boroughs of Staten Island, Queens, and the Bronx at all; temporally, it overrepresents data from Thursday nights. Phase 2, which constitutes the majority of the dataset, consists of 20,816,019 images sampled after October 4 2020, and is much more geographically and temporally representative: it is sampled at all times of the day, on all days of the week, and also covers the entire geographic area of New York City.

Because Phase 2 is much more representative than Phase 1, we conduct our primary analysis of disparities using only data from Phase 2. We additionally conduct numerous validations and bias corrections, described in §4.1, to compensate for non-representative sampling in the dataset. Geographic and temporal coverage during the Phase 2 period is very good. Specifically, 100% of hours during the Phase 2 period are covered; 99.6% of Census Block Groups (CBGs)³ have at least one image, with a mean of 168.2 images per CBG; 88% of roads contained within the borders of New York City are covered by at least one image, using data from OSMNX [6]. Figure 1 summarizes geographic data availability; Figure S1 summarizes temporal data availability.

3.3 Ethical considerations

Our use of this dataset was deemed not human subjects research by our university's Institutional Review Board. Nonetheless, we

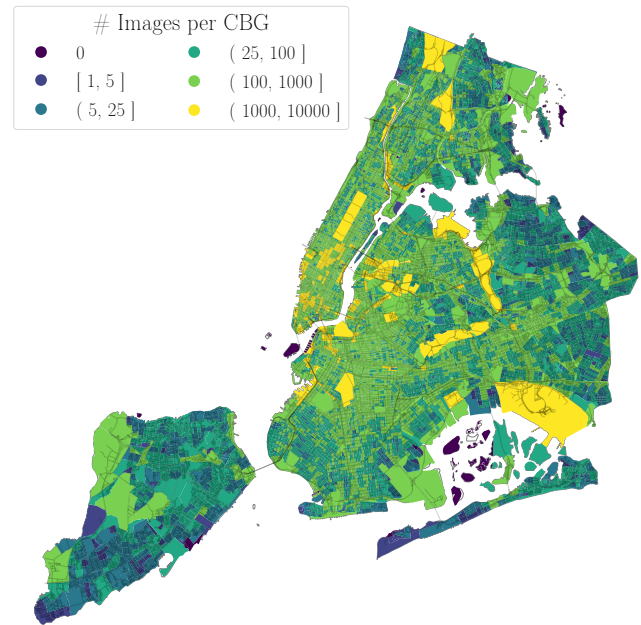


Figure 1: Number of images per Census Block Group in the Phase 2 dataset we use in our primary analysis of disparities.

carefully considered the costs and benefits of this research prior to embarking on it. A key concern for this type of research is that the "subjects"—police vehicles, other vehicles, and other bystanders—do not consent to the initial data collection nor its subsequent use in this research. However, under American Psychological Association Ethics guidelines [3], informed consent procedures can be waived when permitted by law, federal and institutional regulations, or when the research would not reasonably be expected to distress or harm participants and involves naturalistic observations or archival research for which disclosure of responses would not place participants at risk of criminal or civil liability or damage their financial standing, employability or reputation, and for which confidentiality is protected. These provisions are the primary reason why naturalistic observation in public settings does not normally require informed consent [32]. Ultimately, we felt that the potential benefits of this work — namely, quantifying disparities in police deployments to enable more equitable policing — more than offset the costs. This conclusion is consistent with the numerous research works [12, 13, 26, 35, 63, 79] which have used similar datasets capturing public street scenes for research. However, we do believe that it would benefit the public to have broader awareness that data from instruments such as dashcams are now being aggregated on a large-scale basis.

4 STATISTICAL ANALYSIS

In this section, we describe our procedure for assessing disparities in police deployments. First, we describe the mathematical framework to estimate disparities in police deployments while compensating

²New York City is divided into five areas, referred to as boroughs: Queens, Brooklyn, Manhattan, the Bronx, and Staten Island.

³Census Block Groups are Census areas consisting of a couple hundred to a couple thousand people, and are the finest geographic resolution at which we conduct our analysis. New York City has more than 6,000 Census Block Groups.

for potential data and model biases (§4.1). We then provide detail on the deep learning model, which is a key ingredient to this framework, as it detects police vehicles from dashcam images (§4.2).

4.1 Mathematical framework

Our goal is to assess disparities in police vehicle deployments while mitigating the effects of potential model or dataset biases. Here we introduce our mathematical framework for doing this. Before describing the details of the framework, the high-level intuition is that we are reweighting the Nexar data sample, which is sampled from a non-representative set of locations, so that it matches the locations where different demographic groups actually live. For example, to calculate the police deployment levels that Asian residents of New York City experience, we reweight the original data sample to upsample neighborhoods with larger Asian populations.

More formally, let A denote the demographic variable — for example, racial group — over which we wish to assess disparities in police deployment, and let a denote specific values of this variable. Our goal is to estimate the probability a person of group a has a police vehicle visible when they are on a street in their home Census area: in other words, to estimate $\Pr(y = 1|A = a)$, where y indicates whether there is a police vehicle visible. Letting C denote Census area, and summing over Census areas, we can write:

$$\begin{aligned} \Pr(y = 1|A = a) &= \sum_c \Pr(C = c|A = a) \cdot \Pr(y = 1|C = c, A = a) \\ &= \sum_c \Pr(C = c|A = a) \cdot \Pr(y = 1|C = c) \end{aligned} \quad (1)$$

where the second line reflects our assumption that conditional on Census area, the probability a police vehicle is visible is constant across groups: that is, one group within a Census area is no more likely to observe a police vehicle than another. We use fine-grained Census areas throughout our analysis — Census Block Groups — to render this assumption reasonable.⁴

We estimate the first term in the product, $\Pr(C = c|A = a)$, as the fraction of people of group a who live in Census area c . Let N_{ca} denote the number of people of group a living in Census area c . Then, letting $\hat{\Pr}(C = c|A = a)$ denote our estimate of the true probability $\Pr(C = c|A = a)$, our estimator is:

$$\hat{\Pr}(C = c|A = a) = \frac{N_{ca}}{\sum_c N_{ca}} \quad (2)$$

The other term we must estimate to compute Equation 1 is $\Pr(y = 1|C = c)$. Let $\hat{y} = 1$ if our police vehicle detection model detects a police vehicle in an image and $\hat{y} = 0$ otherwise. Then our estimator is:

$$\begin{aligned} \hat{\Pr}(y = 1|C = c) &= \hat{\Pr}(\hat{y} = 1|C = c) \cdot \hat{\Pr}(y = 1|\hat{y} = 1) + \\ &\quad \hat{\Pr}(\hat{y} = 0|C = c) \cdot \hat{\Pr}(y = 1|\hat{y} = 0) \end{aligned} \quad (3)$$

In other words, we estimate the fraction of images which truly have police vehicles in a Census area by taking the fraction classified positive multiplied by the estimated classifier precision $\hat{\Pr}(y = 1|\hat{y} = 1)$, plus the fraction classified negative multiplied by the estimated classifier false omission rate $\hat{\Pr}(y = 1|\hat{y} = 0)$. This procedure compensates for imperfect classifier performance. This procedure will be invalid if classifier performance — i.e., $\Pr(y = 1|\hat{y} = 1)$ and

$\Pr(y = 1|\hat{y} = 0)$ — does not remain constant across Census areas, and in particular if it varies by demographic group a . Given the extensive literature illustrating that classifier performance can vary by by demographic group [7, 11, 14, 38, 76], this is important to verify, and in §4.2.3 we do so.

Overall, our estimation procedure compensates for two types of potential bias. Equation 2 compensates for a *data bias*, reweighting the Nexar dataset (which is sampled from a set of locations which does not necessarily match the population distribution; Figure 1) to match the population distribution of demographic subgroups. This is conceptually similar to inverse propensity weighting procedures [4] which are used to compensate for non-representative data in other settings. Equation 3 compensates for imperfect model performance, and allows us to check that model performance is unbiased (i.e., calibrated) across demographic subgroups.

We note that there are other biases which could prevent us from perfectly estimating disparities in police deployments. In the Discussion section, we discuss two potential biases in more detail. First, the sample of Nexar images we have *within* each Census area c may be biased, in the sense that it may not capture the true probability a resident of that Census area will see a police vehicle at any given time when they are out on the road. For example, if vehicles are prohibited from driving near protest areas, which also have larger police presences, we will not have images of large police presences near protests. It is not possible to correct for this bias with the data we have because 1) the true distribution may differ from the Nexar sampling distribution along unobservable dimensions which we cannot reweight along and 2) we may simply have no Nexar images in some regions of the true distribution (e.g. if all vehicles are banned near protests). A second potential bias is that police vehicles represent only a subset of overall police activity: for example, they do not capture officers on foot. We return to both these points below.

4.2 Deep learning model

We train a deep learning model to identify police vehicles with high accuracy. Our model training pipeline consists of three steps. First, we annotate a large dataset of images for presence of police vehicles (§4.2.1); second, we train a police vehicle detection model on this dataset (§4.2.2); finally, we use a held-out dataset to verify that the model achieves high accuracy on the dataset overall and is not biased with respect to demographic groups (§4.2.3).

4.2.1 Data annotation. We begin by creating training, validation, and test datasets for our police vehicle detection model. Following standard machine learning practice, we use the training set to optimize the parameters of each machine learning model; the validation set to select the model which yields the best performance; and the test dataset to measure the performance of the chosen model. We draw all training images from a random 10% of the dataset, and reserve the remaining 90% for validating the model and analyzing disparities.

A challenge in creating the training set is that the dataset is imbalanced — the vast majority of images do not contain police vehicles — and highly imbalanced datasets can produce inferior model performance [33]. To mitigate this issue, we construct a more balanced training set by sampling images near police stations,

⁴We use data from the 2020 American Community Survey.

which are more likely to contain police vehicles: specifically, we filter for images within 0.25 miles of one of NYC’s 77 police precinct stations, which raises the proportion of images with police vehicles from about 1% to roughly 7%. The risk of this training strategy is that it potentially introduces distribution shift [39], which can also harm model performance, because images near police stations may differ from the overall image distribution; below, we describe how we check for this concern.

We collect annotations for 15,250 images using Scale.AI [61], an outsourcing platform for annotations similar to Amazon Mechanical Turk. Each image is annotated with rectangular bounding boxes to identify police vehicles. We utilize evaluation tasks (images pre-verified with ground truth annotations by us) on the platform to disqualify inaccurate labelers. Every image labeled by a Scale.AI annotator is then submitted to a review phase, where a separate, high-scoring labeler checks each image for accuracy. In choosing which types of NYPD vehicles to have annotators label, we account for the fact that the NYPD has several divisions, including the prisoner transport, building maintenance, and museum units [55], that are not involved in active policing. We instruct annotators to only label NYPD sedans, SUVs, compacts, and trucks (but not buses or 4x4 trucks) to mitigate this. In total, 1,088 images from Scale.AI have police vehicles; we add these to the training set.

To further increase the size of the training dataset, we use the model trained on the Scale.AI annotations to select images which it predicts have police vehicles (at a confidence threshold of 0.9). Specifically, we apply the model to 450,000 images selected at random from the 10% training sample (not just images near police stations); examine all images passing the 0.9 confidence threshold and manually verify whether these images do, in fact, contain police vehicles, and add the images with verified labels to the train set. The final training set consists of 4,748 images with police vehicles and 4,701 images without police vehicles.

In contrast to the training set, which oversamples images near police stations, the validation and test splits are both random samples of 20,000 images from the overall dataset. A random sample is essential because it provides us with an unbiased measure of performance, which our estimation framework requires (§4.1). To annotate the validation and test sets, we use Scale.AI to collect annotations and manually audit all images labeled as positive for correctness. The validation set includes 247 positive images (i.e., which contain police vehicles) out of 20,000, and the test set includes 239 positive images out of 20,000.

4.2.2 Model training. We frame our police vehicle detection problem as a rectangular object detection problem. We fine-tune models from the widely-used You Only Look Once (YOLO) suite of object detection models [74] which have been pretrained on the MS-COCO dataset to detect object classes including vehicles [24]. We compare performance from a range of model configurations (for example, we compare performance from model variants YOLO5X [70] and YOLO7-E6E [74], and experiment with hyperparameter evolution). We use average precision on the validation set to select the model which yields optimal performance.⁵ Our final model is trained using

⁵Because each image can potentially contain multiple police vehicles and multiple model predictions, to compute average precision, we define an overall model prediction for each image as the maximum model prediction for any object detection (or as 0 if

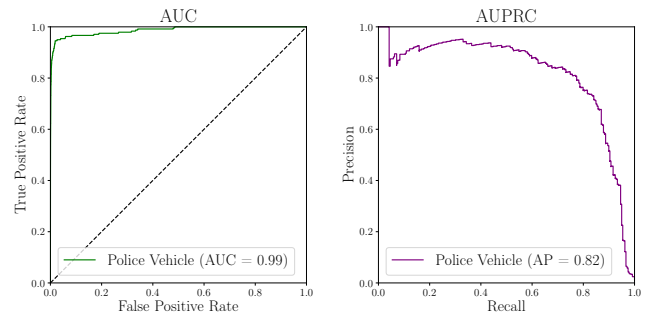


Figure 2: Receiver Operating Characteristic (ROC) and Precision-Recall curves for our final model, evaluated on the test set.

the default hyperparameter configuration for YOLOv7. To choose a threshold for positive classification, we use the threshold (0.77) which maximizes the F-score on the validation set.

4.2.3 Model validation. Using the held-out test set, we first assess model performance both on the dataset as a whole, and then verify that the model achieves consistent performance across subsets of the dataset. Figure 2 shows the precision-recall curve and ROC curve for the test set, and Table S2 reports average precision, AUC (i.e., AUROC), precision and recall. The model achieves high performance (test set AUC 0.99, average precision 0.82). This speaks to one benefit of our methodology: because we have chosen a relatively feasible object detection task (likelier easier than identifying surveillance cameras, as attempted in previous work [63], which are much smaller) we are able to achieve high performance. Figure 3 shows specific true positive, false positive, and false negative classifications on the dataset, demonstrating that the model can recognize police vehicles in a wide variety of images, but also that it can be confused by optical illusions and poor light conditions.

Our estimation framework (§4.1) requires that our model is *calibrated* across subgroups — that is, $\Pr(y = 1|\hat{y} = 1)$ and $\Pr(y = 1|\hat{y} = 0)$ remain similar across subgroups. We therefore assess whether this is the case using the test set, and as an additional validation assess model average precision and AUC across subgroups. Figure 4 shows that the model is calibrated across demographic variables (e.g. % White, % Black, population density, median household income, and Manhattan vs. non-Manhattan). The model is also calibrated across a number of other variables (whether the image is taken during the day, on the weekend, or during phase 1). Table S1 shows that AUC and average precision remain consistently high across subgroups (AUC 0.98–0.99 for all subgroups; average precision 0.77–0.84). The only evidence of miscalibration we see does not occur in the variables over which we assess disparities, and thus does not threaten the validity of those estimates. Specifically, the model is somewhat miscalibrated by distance from nearest police stations: $\Pr(y = 1|\hat{y} = 0)$ is higher for images closer to police stations. This

no objects are detected). We define the true label for each image as 1 if it contains any police vehicles, and 0 otherwise. We do this because it allows us to compute precision and recall statistics at the per-image level, not the per-object level, and our primary analysis is at the per-image level.



Figure 3: Examples of police vehicle detections. Notice the different times of day, differing camera perspectives, and multiple types of vehicles. Initially, the object detection model was easily confused by NYC taxis, buses, white sedans, and white SUVs. As more training data was added, the model was more difficult to confuse, with incorrect classifications occurring typically with optical illusions and poor light conditions. Note that images are cropped and zoomed to better depict annotated regions; raw images are 1280x720.

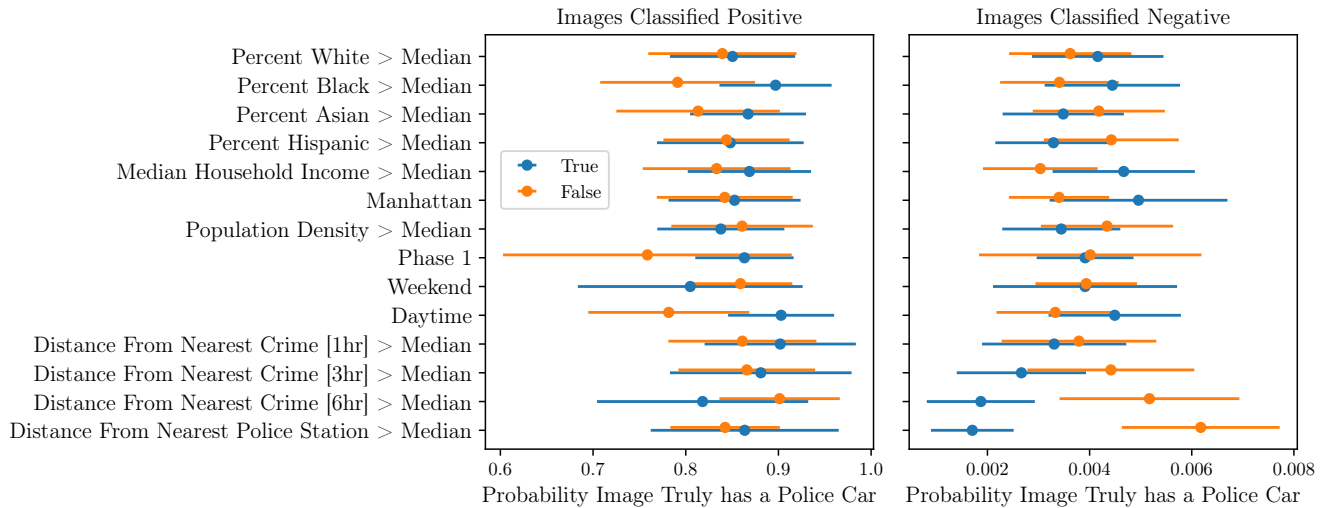


Figure 4: Calibration plots by subgroup, evaluated on the test set. The left plot shows $\Pr(y = 1 | \hat{y} = 1)$; the right subfigure shows $\Pr(y = 1 | \hat{y} = 0)$. Lines indicate uncertainty in estimates, calculated as 1.96 times the Bernoulli standard error; if the difference between the orange and blue dots is smaller than the uncertainty, it indicates that there is no statistically significant difference in model performance across the stratification.

is unsurprising, given the overdensity of police vehicles near police stations, and may also result from the distribution shift from training set to validation/testing set. We experiment with two fixes for this issue: first, we try upsampling images which the model misclassifies, but find it harms overall performance significantly. Second, we recalibrate the final model using distance from nearest

police station as a covariate. This recalibration does not significantly affect our main results, as expected, so we do not use the recalibrated model for simplicity. There is also slight evidence of miscalibration by distance from nearest crime, but it is inconsistent and depends on the temporal threshold used for determining the nearest crime.

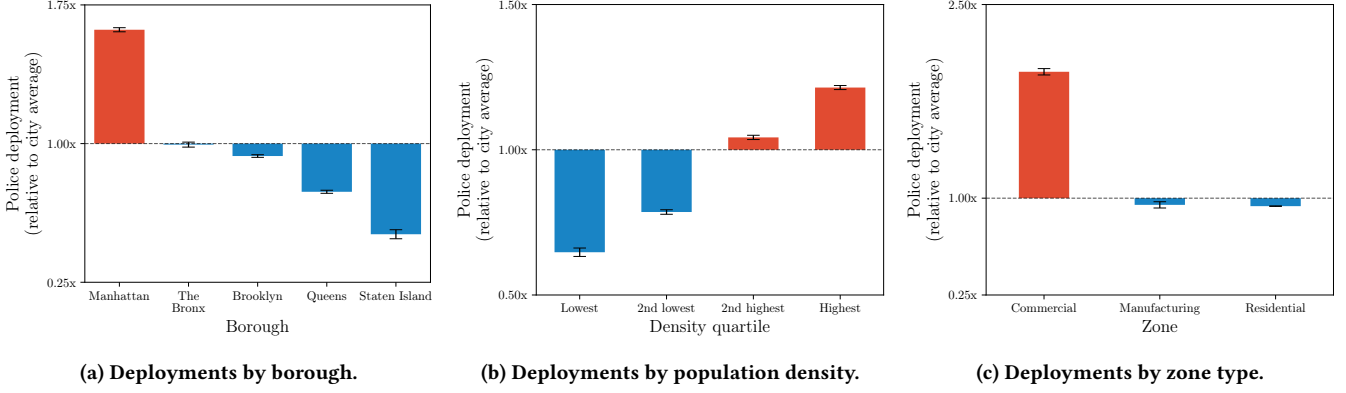


Figure 5: Disparities in police deployments by borough, density, and zone.

Overall, the model is generally well-calibrated and achieves high performance across subgroups over which we assess disparities, as our mathematical framework requires. At our positive classification threshold of 0.77, the model classifies 233,596 images (out of 24,803,854 total) as containing a police vehicle. From the validation set, we estimate $\widehat{\Pr}(y = 1|\hat{y} = 1)$ and $\widehat{\Pr}(y = 1|\hat{y} = 0)$, and use them as described in §4.1 to compute disparities in police deployments across demographic groups.

As an additional validation, we assess whether the police detections correlate as expected with external data. Specifically, we assess whether Census Block Groups whose images have a lower mean distance to crime or to police stations are likelier to have police vehicles detected, as one would expect. We find that both correlations hold ($p < 0.001$, Pearson correlation across Census Block Groups). This provides an additional validation beyond assessing model performance, since it suggests that police vehicle presence shows the same correlations with crime and police stations which we would expect police activity as a whole to display.

5 ANALYSIS OF DISPARITIES

We examine disparities in police deployments — i.e., $\Pr(y = 1|A = a)$, computed as described in §4.1 — by racial group (white, Black, Hispanic, and Asian), zone type (commercial, residential, and manufacturing), population density, median household income, borough, and neighborhood. Throughout the results, we express police deployments relative to the population-weighted city overall average: e.g., $1.6\times$ for a group indicates that deployment levels are 1.6 times the overall average. We compute error bars for all estimates by bootstrapping; reported errors are 1.96 times the standard deviation across bootstrapped datasets.

We find significant disparities in police deployments across boroughs (Figure 5) and neighborhoods (Figure 6). The borough with the highest police deployments (Manhattan) has levels $1.62 \pm 0.01\times$ the city average; the borough with the lowest police deployments (Staten Island) has levels $0.51 \pm 0.02\times$ the city average, a more than 3-fold difference. Similarly, the neighborhood with the highest police deployment levels (Gramercy) has levels $4.28 \pm 0.10\times$ the city average; the neighborhood with the lowest police deployments (Arden Heights-Rossville) has levels $0.23 \pm 0.02\times$ the city average,

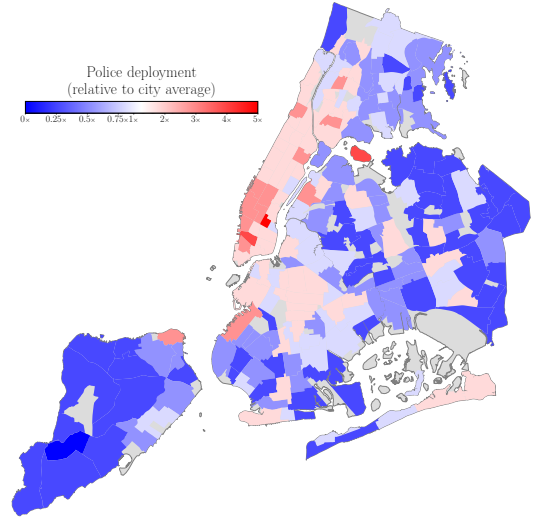


Figure 6: Map of police deployment throughout NYC, expressed relative to the city average. Grey areas are those with zero population in Census data, including airports, cemeteries, and parks.

a nearly 20-fold difference. Examining the neighborhoods with the highest police levels reveals a striking mixture of two types of neighborhoods: some of the wealthiest areas in New York (Hudson Yards, once described as a “billionaire’s fantasy city” [18]) and some of the most heavily policed (Rikers Island, which houses a prison complex).

High police deployment levels in some neighborhoods may be driven by idiosyncratic neighborhood attributes — for example, Gramercy is home to a major police training center on two busy thoroughfares. To more systematically examine spatial disparities,

therefore, we examine correlations with zone and density (Figure 5). Commercial zones have much higher police deployments ($1.98 \pm 0.03\times$ the city average) than residential or manufacturing zones. Police deployments also increase dramatically with population density, from $0.65 \pm 0.01\times$ the city average in areas in the lowest quartile of density to $1.21 \pm 0.01\times$ the city average in areas in the highest quartile of density.⁶ Overall, this analysis suggests that one place high police levels are observed is dense, commercial regions of downtown Manhattan.

We next analyze disparities in police deployments by race and socioeconomic status (Figure 7a), restricting the analysis to residential zones to capture where people actually live. We observe substantial racial disparities. Black and Hispanic residents face police deployment levels $1.08 \pm 0.01\times$ the city average, as compared to $0.95 \pm 0.004\times$ for white residents and $0.81 \pm 0.01\times$ for Asian residents, a 34% discrepancy between the highest and lowest race groups. This is consistent with previously observed disparities in policing in New York City [27, 36, 52] where Black and Hispanic residents were more heavily policed. (In Figure S2 we plot racial disparities without restricting to residential zones, yielding consistent results — Black and Hispanic residents face higher police deployments than white and Hispanic residents — although the disparities narrow somewhat, to a 19% gap between the highest and lowest race groups.)

We also observe disparities by Census Block Group (CBG) median income (Figure 7b). Again restricting to residential zones, police deployments increase as income decreases: residents of CBGs in the lowest median household income quartile face police deployments which are $1.19 \pm 0.01\times$ the city average, while residents of CBGs in the highest median household income quartile face police deployments only $0.91 \pm 0.01\times$ the city average. Interestingly, the relationship between income and deployment becomes U-shaped if we do not restrict to residential zones (Figure S2): neighborhoods in the *top* quartile of median income, and the *bottom* quartile of median income, both experience higher police deployments than neighborhoods in the middle two quartiles. In other words, top-quartile income areas go from having the lowest police deployments (when we analyze only residential zones) to the highest (when we analyze all areas). The difference likely occurs because of higher-income commercial zones which are heavily policed.

We note that we report racial and socioeconomic disparities without attempting to control for other covariates for two reasons. First, if New York City residents of different races face different levels of police deployments, that disparate impact is itself important; it can also bias algorithms trained on downstream policing data irrespective of the true causal mechanism. Second, controlling for other factors in policing data can be difficult to interpret, introducing concerns about omitted variable bias and model misspecification which make it difficult to identify which factor is truly the “cause” of higher police deployments. For the sake of transparency and simplicity, therefore, we report results by stratifying each variable separately, noting that these disparities are themselves important but that multiple causal mechanisms may underlie them.

⁶We compute quartile boundaries at the image level.

6 DISCUSSION

We present a novel methodology for studying disparities in police deployments. We show that a deep learning model can be applied to tens of millions of dashcam images to identify police vehicles with high accuracy, and introduce a principled framework for estimating disparities in a way that mitigates data and model biases. We find substantial spatial inequality in where police vehicles are deployed. Residential areas with higher proportions of Black and Hispanic residents, or lower-income residents, have significantly higher police deployments. But so, too, do dense, wealthy commercial areas in downtown Manhattan like Hudson Yards; this finding that police deployments can be high in wealthy neighborhoods is consistent with past work showing that gentrification correlates with increased police presence [5]. Overall, our analysis speaks to a complex, multifaceted picture of police deployment in New York.

We employ a number of techniques to reduce bias in our pipeline, providing a template for future computational social science research on non-representative dashcam data. However, two unavoidable limitations are important to acknowledge. First, while we reweight the sample such that each Census area is weighted proportional to its population, it is possible that the sample of images we have *within each Census area* is biased. In other words, the police vehicle frequency within our Nexar data sample for each Census area may not accurately represent the true police vehicle frequency within that Census area. News articles document, for example, how protests after George Floyd’s death blocked traffic throughout New York City [23], which would obviously preclude getting dashcam images from some areas; it is also plausible that ride-share drivers would seek to avoid areas with protest activity. Since protest activity likely correlates both with our data availability and with police vehicle presence, this could plausibly bias our estimates. We mitigate the specific concern about protests by conducting our analysis of disparities using data only from October 2020 onward, when protest activity was reduced. Nonetheless, the sampling process for the Nexar dataset is somewhat opaque and the time period analyzed was full of dramatic changes in social distancing policy and responses to policing so it would be unwise to assume that no other biases remain. Similarly, it is possible the time period sampled captures trends specific to the COVID-19 pandemic. At the same time, that time period itself is of interest because it is known that there were racial disparities in policing through the pandemic. Kajeepeta et al. [37] find that from March to May 2020, a one standard deviation increase in percentage of Black residents was associated with a 73% increase in the COVID-19-specific summons rate and a 34% increase in the public health and nuisance arrest rate. Finally, we assume that people in the same Census area have the same probability of seeing a police vehicle irrespective of their group; while the Census areas we analyze (Census Block Groups) are quite small, rendering this assumption more plausible, it is possible there is additional variation correlated with group within Census areas.

A second important limitation is that identifiable police vehicles are only a partial proxy for all policing activity. They do not capture, for example, officers on foot or unmarked vehicles. We did attempt to train a computer vision model to identify officers on foot; however, we found that the annotations we obtained were considerably lower quality than annotations for police vehicles,

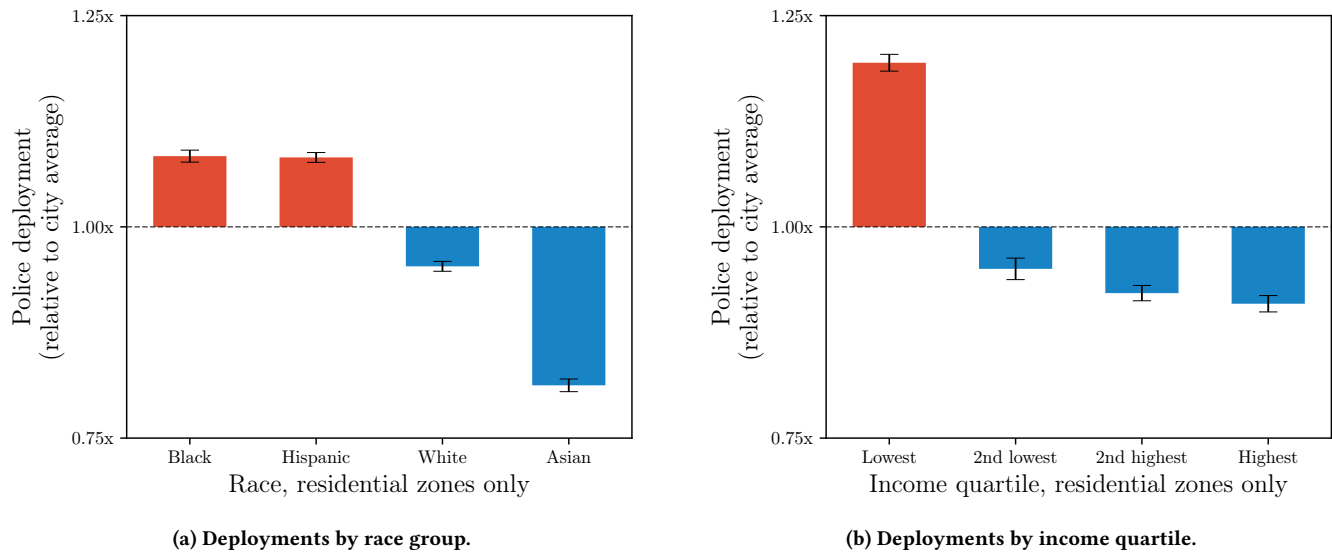


Figure 7: Disparities in police deployments by race group (a) and median household income (b). Plots are made using data from Census Block Groups in residential zones. Figure S2 shows the corresponding plots without filtering for residential zones; results for race are qualitatively similar, but results for income show a U-shaped trend, for reasons we discuss in the main text.

with annotators frequently confusing police officers with construction workers, cyclists with reflective vests, and other individuals in dark clothing. Given this, we opted to focus our analysis on police vehicles to ensure we could train a classifier with high accuracy. The high performance of the classifier likely stems in part from the iconic and unique branding of NYPD vehicles, which makes them very easy to spot and label in training images. While it is important to be aware that our analysis captures only one dimension of police activity, police vehicles are, themselves, an important indicator of policing activity: for example, they carry out traffic stops, which are one of the most common ways that Americans interact with police [53].

We assert our approach and results constitute a key step towards transparency in policing, proving it is possible to derive deployment patterns from dashcam data. Our results have several implications for equity and algorithmic fairness. First, our work reveals substantial inequality in police vehicle deployments. Police vehicles are more heavily deployed in neighborhoods with larger Black and Hispanic populations, recapitulating historic patterns of inequality in New York policing [27, 28, 52]. These disparities are themselves of concern, and could also propagate into algorithms which made use of data from outcomes downstream of deployments, including arrests [15, 16, 20, 42, 45, 54, 58, 62].

Second, our work highlights the potential of dashcam data for auditing government agencies for efficiency and equity, including but not limited to the police, as well as for computational social science more broadly. For example, similar methodology could also be used to identify double parking issues, dynamic obstacles to accessibility, or weather-induced hazards, and then to study disparities in government service response times, following previous work [41].

Dashcam data is becoming increasingly available [47, 51], and dashcam is being postulated as a medium for future work [43, 44]. Our methods could easily be extended to other cities or detection tasks.

Finally, our work serves as a call for the police themselves to release better deployment data. In the past police have resisted this [48] on the grounds that it would undermine public safety. But this claim stretches plausibility on multiple grounds: even *aggregated* deployment data would be very useful for detecting disparities, as we show and past work has also found [49, 50]. And far more granular data on many measures of policing activity is already publicly available — including stops, searches, arrests, citations, and shootings [34, 53] — making it unclear how aggregated deployment data uniquely compromises police operations. The benefits would be substantial: public release of aggregated deployment data would democratize the ability to audit the police. This project was performed using highly specialized resources — terabytes of data, a customized annotation pipeline and computer hardware, and hundreds of hours of model training — and even then incurred unavoidable biases. This pipeline is not, in short, a realistic way for the average citizen to monitor the police. *Quis custodiet ipsos custodes?* The answer cannot be “only those with 25 million dashcam images.”

Code availability: Code used in this analysis is available at this GitHub repository: <https://github.com/mattwfranchi/police-deployment-patterns>.

Dataset availability: It is our intent to share the data from this paper with other researchers, with protections to prevent abuse. To learn how to access to data from the study, see <https://github.com/mattwfranchi/police-deployment-patterns/wiki/Dataset-Access>.

Acknowledgments: The authors thank Serina Chang, Nikhil Garg, Nakia Kenon, Hao Sheng, and Maria Teresa Parreira for helpful

comments, and Gabriel Agostini and Ilan Mandel for assistance with data and analysis. The Nexar dashcam dataset used in this study was licensed for use as part of NSF IIS-2028009 RAPID: Tracking Urban Mobility and Occupancy under Social Distancing Policy. WJ was supported by an Amazon Research Award. EP was supported by a Google Research Scholar award, an NSF CAREER award, a CIFAR Azrieli Global scholarship, a LinkedIn Research Award, and a Future Fund Regrant. JZ was partially supported by the United States Air Force and DARPA under contracts FA8750-20-C-0156, FA8750-20-C-0074, and FA8750-20-C0155 (SDCPS Program).

REFERENCES

- [1] Desmond Ang. 2021. The effects of police violence on inner-city students. *The Quarterly Journal of Economics* 136, 1 (2021), 115–168.
- [2] Shamena Anwar and Hanming Fang. 2006. An alternative test of racial prejudice in motor vehicle searches: Theory and evidence. *American Economic Review* 96, 1 (2006), 127–151.
- [3] American Psychological Association et al. 2002. Ethical principles of psychologists and code of conduct. *American psychologist* 57, 12 (2002), 1060–1073.
- [4] Peter C Austin and Elizabeth A Stuart. 2015. Moving towards best practice when using inverse probability of treatment weighting (IPTW) using the propensity score to estimate causal treatment effects in observational studies. *Statistics in medicine* 34, 28 (2015), 3661–3679.
- [5] Brenden Beck. 2020. Policing gentrification: Stops and low-level arrests during demographic change and real estate reinvestment. *City & Community* 19, 1 (2020), 245–272.
- [6] Geoff Boeing. 2017. OSMnx: New Methods for Acquiring, Constructing, Analyzing, and Visualizing Complex Street Networks. *Computers, Environment and Urban Systems* 65 (Sept. 2017), 126–139. <https://doi.org/10.1016/j.compenvurbsys.2017.05.004> arXiv:1611.01890 [physics].
- [7] 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.
- [8] Andrew Campbell, Alan Both, and Qian Chayn Sun. 2019. Detecting and mapping traffic signs from Google Street View images using deep learning and GIS. *Computers, Environment and Urban Systems* 77 (2019), 101350.
- [9] Cindy Chang and Ben Poston. 2019. LAPD will drastically cut back on pulling over random vehicles over racial bias concerns. *Los Angeles Times* (2019).
- [10] Cindy Chang and Ben Poston. 2019. Stop-and-frisk in a car: Elite LAPD unit disproportionately stopped black drivers, data show. *Los Angeles Times* (2019).
- [11] Irene Y Chen, Emma Pierson, Sherri Rose, Shalmali Joshi, Kadija Ferryman, and Marzyeh Ghassemi. 2021. Ethical machine learning in healthcare. *Annual review of biomedical data science* 4 (2021), 123–144.
- [12] Tahiya Chowdhury, Ansh Bhatti, Ilan Mandel, Taqiya Ehsan, Wendy Ju, and Jorge Ortiz. 2021. Towards sensing urban-scale COVID-19 policy compliance in New York City. In *Proceedings of the 8th ACM International Conference on Systems for Energy-Efficient Buildings, Cities, and Transportation*. 353–356.
- [13] Tahiya Chowdhury, Qizhen Ding, Ilan Mandel, Wendy Ju, and Jorge Ortiz. 2021. Tracking urban heartbeat and policy compliance through vision and language-based sensing. In *Proceedings of the 8th ACM International Conference on Systems for Energy-Efficient Buildings, Cities, and Transportation*. 302–306.
- [14] Sam Corbett-Davies and Sharad Goel. 2018. The measure and mismeasure of fairness: A critical review of fair machine learning. *arXiv preprint arXiv:1808.00023* (2018).
- [15] Sam Corbett-Davies, Emma Pierson, Avi Feller, and Sharad Goel. 2016. A computer program used for bail and sentencing decisions was labeled biased against blacks. It's actually not that clear. *Washington Post* 17 (2016).
- [16] Sam Corbett-Davies, Emma Pierson, Avi Feller, Sharad Goel, and Aziz Huq. 2017. Algorithmic decision making and the cost of fairness. In *Proceedings of the 23rd acm sigkdd international conference on knowledge discovery and data mining*. 797–806.
- [17] Bahar Dadashova, Chiara Silvestri Dobrovolny, and Mahmood Tabesh. 2021. Detecting Pavement Distresses Using Crowdsourced Dashcam Camera Images. (2021).
- [18] Justin Davidson. 2019. Hudson Yards is a billionaire's fantasy city and you never have to leave — provided you can pay for it. *New York Magazine* (2019).
- [19] Elizabeth Davis, Anthony Whyde, and Lynn Langton. 2018. Contacts between police and the public, 2015. *US Department of Justice Office of Justice Programs Bureau of Justice Statistics Special Report* (2018), 1–33.
- [20] Angel Diaz. 2021. Data-driven policing's threat to our constitutional rights. *Brookings Institution* (2021).
- [21] Frank Edwards, Hedwig Lee, and Michael Esposito. 2019. Risk of being killed by police use of force in the United States by age, race-ethnicity, and sex. *Proceedings of the national academy of sciences* 116, 34 (2019), 16793–16798.
- [22] Charles R Epp, Steven Maynard-Moody, and Donald P Haider-Markel. 2014. *Pulled over: How police stops define race and citizenship*. University of Chicago Press.
- [23] Alan Feuer and Azi Paybarah. 2020. Thousands Protest in N.Y.C., Clashing With Police Across All 5 Boroughs. *The New York Times* (2020).
- [24] David Fleet, Tomas Pajdla, Bernt Schiele, and Tinne Tuytelaars (Eds.). 2014. *Computer Vision – ECCV 2014: 13th European Conference, Zurich, Switzerland, September 6–12, 2014, Proceedings, Part V*. Lecture Notes in Computer Science, Vol. 8693. Springer International Publishing, Cham. <https://doi.org/10.1007/978-3-319-10602-1>
- [25] Roland G Fryer Jr. 2019. An empirical analysis of racial differences in police use of force. *Journal of Political Economy* 127, 3 (2019), 1210–1261.
- [26] Timnit Gebru, Jonathan Krause, Yilun Wang, Duyun Chen, Jia Deng, Erez Lieberman Aiden, and Li Fei-Fei. 2017. Using deep learning and Google Street View to estimate the demographic makeup of neighborhoods across the United States. *Proceedings of the National Academy of Sciences* 114, 50 (2017), 13108–13113.
- [27] Andrew Gelman, Jeffrey Fagan, and Alex Kiss. 2007. An analysis of the New York City police department's "stop-and-frisk" policy in the context of claims of racial bias. *Journal of the American statistical association* 102, 479 (2007), 813–823.
- [28] Sharad Goel, Justin M Rao, and Ravi Shroff. 2016. Precinct or prejudice? Understanding racial disparities in New York City's stop-and-frisk policy. *The Annals of Applied Statistics* 10, 1 (2016), 365–394.
- [29] Felipe Goncalves and Steven Mello. 2021. A few bad apples? Racial bias in policing. *American Economic Review* 111, 5 (2021), 1406–41.
- [30] Jeffrey Grogger and Greg Ridgeway. 2006. Testing for racial profiling in traffic stops from behind a veil of darkness. *J. Amer. Statist. Assoc.* 101, 475 (2006), 878–887.
- [31] Andrea Marie Headley, Stewart J D'Alessio, and Lisa Stolzenberg. 2020. The effect of a complainant's race and ethnicity on dispositional outcome in police misconduct cases in Chicago. *Race and justice* 10, 1 (2020), 43–61.
- [32] Elizabeth M Hill and Mary N Monahan. 2011. Naturalistic Observation in Public Settings: Applying for Institutional Review Board Approval. *Human Ethology Bulletin* (2011).
- [33] Nathalie Japkowicz and Shaju Stephen. 2002. The class imbalance problem: A systematic study. *Intelligent data analysis* 6, 5 (2002), 429–449.
- [34] Jennifer Jenkins, Monika Mathur, John Muyskens, Razzan Nakhlawi, Steven Rich, and Andrew Ba Tran. 2023. Police shootings database 2015–2023s. *The Washington Post* (2023). <https://www.washingtonpost.com/graphics/investigations/police-shootings-database/>
- [35] Wendy Ju, Sharon Yavo-Ayalon, Ilan Mandel, Federico Saldarini, Natalie Friedman, Srinath Sibi, JD Zamfirescu-Pereira, and Jorge Ortiz. 2020. Tracking urban mobility and occupancy under social distancing policy. *Digital Government: Research and Practice* 1, 4 (2020), 1–12.
- [36] Jongbin Jung, Sam Corbett-Davies, Ravi Shroff, and Sharad Goel. 2018. Omitted and included variable bias in tests for disparate impact. *arXiv preprint arXiv:1809.05651* (2018).
- [37] Sandhya Kajeepeta, Emilie Bruzelius, Jessica Z. Ho, and Seth J. Prins. 2022. Policing the pandemic: estimating spatial and racialized inequities in New York City police enforcement of COVID-19 mandates. *Critical Public Health* 32, 1 (Jan. 2022), 56–67. <https://doi.org/10.1080/09581596.2021.1987387>
- [38] Allison Koenecke, Andrew Nam, Emily Lake, Joe Nudell, Minnie Quartey, Zion Mengesha, Connor Toups, John R Rickford, Dan Jurafsky, and Sharad Goel. 2020. Racial disparities in automated speech recognition. *Proceedings of the National Academy of Sciences* 117, 14 (2020), 7684–7689.
- [39] Pang Wei Koh, Shiori Sagawa, Henrik Marklund, Sang Michael Xie, Marvin Zhang, Akshay Balsubramani, Weihua Hu, Michihiro Yasunaga, Richard Lanus Phillips, Irena Gao, et al. 2021. Wilds: A benchmark of in-the-wild distribution shifts. In *International Conference on Machine Learning*. PMLR, 5637–5664.
- [40] Vladimir A Krylov, Eamonn Kenny, and Rozenn Dahyot. 2018. Automatic discovery and geotagging of objects from street view imagery. *Remote Sensing* 10, 5 (2018), 661.
- [41] Benjamin Laufer, Emma Pierson, and Nikhil Garg. Working paper, 2023; ICML Workshop on Responsible Decision Making in Dynamic Environments, 2022. Detecting Disparities in Capacity-Constrained Service Allocations. (Working paper, 2023; ICML Workshop on Responsible Decision Making in Dynamic Environments, 2022).
- [42] Laura and John Arnold Foundation. 2016. Public safety assessment: Risk factors and formula.
- [43] Xiaojiang Li. 2021. Examining the spatial distribution and temporal change of the green view index in New York City using Google Street View images and deep learning. *Environment and Planning B: Urban Analytics and City Science* 48, 7 (Sept. 2021), 2039–2054. <https://doi.org/10.1177/2399808320962511> Publisher: SAGE Publications Ltd STM.
- [44] Xiaojiang Li, Chuanrong Zhang, Weidong Li, Robert Ricard, Qingyan Meng, and Weixing Zhang. 2015. Assessing street-level urban greenery using Google Street View and a modified green view index. *Urban Forestry & Urban Greening* 14, 3 (2015), 675–685. <https://doi.org/10.1016/j.ufug.2015.06.006>

- [45] Kristian Lum and William Isaac. 2016. To predict and serve? *Significance* 13, 5 (2016), 14–19.
- [46] Eddy Maddalena, Luis-Daniel Ibáñez, and Elena Simperl. 2020. Mapping points of interest through street view imagery and paid crowdsourcing. *ACM Transactions on Intelligent Systems and Technology (TIST)* 11, 5 (2020), 1–28.
- [47] Vashisht Madhavan and Trevor Darrell. 2017. *The BDD-Nexar Collective: A Large-Scale, Crowdsourced, Dataset of Driving Scenes*. Master's thesis. EECS Department, University of California, Berkeley. <http://www2.eecs.berkeley.edu/Pubs/TechRpts/2017/EECS-2017-113.html>
- [48] Benjamin Mueller. 2017. As Troopers Are Diverted, Deployment Data Remains Restricted. (08 2017). <https://www.nytimes.com/2017/08/17/nyregion/new-york-state-police-deployment-data.html>
- [49] Benjamin Mueller and Al Baker. 2016. Rift Between Officers and Residents as Killings Persist in South Bronx. *The New York Times* (2016).
- [50] Benjamin Mueller and Al Baker. 2017. New York Police Urged to Fix Inequities in Deployment of Investigators. *The New York Times* (2017).
- [51] Xingchao Peng, Ben Usman, Neela Kaushik, Dequan Wang, Judy Hoffman, and Kate Saenko. 2018. VisDA: A Synthetic-to-Real Benchmark for Visual Domain Adaptation. In *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition Workshops (CVPRW)*. IEEE, Salt Lake City, UT, USA, 2102–21025. <https://doi.org/10.1109/CVPRW.2018.00271>
- [52] Emma Pierson, Sam Corbett-Davies, and Sharad Goel. 2018. Fast threshold tests for detecting discrimination. In *International conference on artificial intelligence and statistics*. PMLR, 96–105.
- [53] Emma Pierson, Camelia Simoiu, Jan Overgoor, Sam Corbett-Davies, Daniel Jensen, Amy Shoemaker, Vignesh Ramachandran, Phoebe Barghouty, Cheryl Phillips, Ravi Shroff, et al. 2020. A large-scale analysis of racial disparities in police stops across the United States. *Nature human behaviour* 4, 7 (2020), 736–745.
- [54] Michael Pinard. 2018. Predicting more biased policing in Baltimore. *The Baltimore Sun* (2018).
- [55] PoliceCarWebsite.net. 2023. NYPD Police Cars. <http://www.policecarwebsite.net/fc/ny/nypd/nypddiv.html>
- [56] Ben Poston and Cindy Chang. 2019. LAPD searches blacks and Latinos more. But they're less likely to have contraband than whites. *The Los Angeles Times* (2019).
- [57] Rajeev Ramchand, Rosalie Luccardo Pacula, and Martin Y Iguchi. 2006. Racial differences in marijuana-users' risk of arrest in the United States. *Drug and alcohol dependence* 84, 3 (2006), 264–272.
- [58] Rashida Richardson, Jason M Schultz, and Kate Crawford. 2019. Dirty data, bad predictions: How civil rights violations impact police data, predictive policing systems, and justice. *NYUL Rev. Online* 94 (2019), 15.
- [59] Greg Ridgeway and John M MacDonald. 2009. Doubly robust internal benchmarking and false discovery rates for detecting racial bias in police stops. *J. Amer. Statist. Assoc.* 104, 486 (2009), 661–668.
- [60] Konstantin Riedl, Sebastian Huber, Maximilian Bömer, Julian Kreibich, Felix Nobis, and Johannes Betz. 2020. Importance of Contextual Information for the Detection of Road Damages. In *2020 Fifteenth International Conference on Ecological Vehicles and Renewable Energies (EVER)*. IEEE, 1–7.
- [61] Scale.AI. 2023. Scale.AI. <https://scale.com>
- [62] Andrew D Selbst. 2017. Disparate impact in big data policing. *Ga. L. Rev.* 52 (2017), 109.
- [63] Hao Sheng, Keniel Yao, and Sharad Goel. 2021. Surveilling Surveillance: Estimating the Prevalence of Surveillance Cameras with Street View Data. In *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*. ACM, Virtual Event USA, 221–230. <https://doi.org/10.1145/3461702.3462525>
- [64] Camelia Simoiu, Sam Corbett-Davies, and Sharad Goel. 2017. The problem of infra-marginality in outcome tests for discrimination. (2017).
- [65] Samuel Sinyangwe et al. 2022. The Police Scorecard. (2022). <https://policescorecard.org/>
- [66] Visual Investigations Team. 2023. Visual investigations of police misconduct. *The New York Times* (2023). <https://www.nytimes.com/spotlight/visual-investigations-police-misconduct>
- [67] *Floyd v. City of New York*. 2013. 959 F. Supp. 2d 540 - Dist. Court, SD New York. (2013).
- [68] Andrew Thirlwell and Ognjen Arandjelović. 2020. Big data driven detection of trees in suburban scenes using visual spectrum eye level photography. *Sensors* 20, 11 (2020), 3051.
- [69] Victor J. D. Tsai, Jyun-Han Chen, Pei-Shan Tsai, Hsun-Sheng Huang, and Ke-Jhong Chen. 2017. Exploring Traffic Features on Google Street View Images. In *Proceedings of the 2017 International Conference on Wireless Communications, Networking and Applications (Shenzhen, China) (WCNA 2017)*. Association for Computing Machinery, New York, NY, USA, 230–234. <https://doi.org/10.1145/3180496.3180638>
- [70] UltraLytics. 2021. YOLOv5. <https://github.com/ultralytics/yolov5/releases>
- [71] United States Department of Justice Civil Rights Division and United States Attorney's Office for Northern District of Illinois. 2017. Investigation of the Chicago Police Department. (2017).
- [72] Vinay Vishnani, Anikait Adhya, Chinmay Bajpai, Priya Chimurkar, and Kumar Khadagle. 2020. Manhole detection using image processing on google street view imagery. In *2020 Third International Conference on Smart Systems and Inventive Technology (ICSSIT)*. IEEE, 684–688.
- [73] Rob Voigt, Nicholas P Camp, Vinodkumar Prabhakaran, William L Hamilton, Rebecca C Hetey, Camilla M Griffiths, David Jurgens, Dan Jurafsky, and Jennifer L Eberhardt. 2017. Language from police body camera footage shows racial disparities in officer respect. *Proceedings of the National Academy of Sciences* 114, 25 (2017), 6521–6526.
- [74] Chien-Yao Wang, Alexey Bochkovskiy, and Hong-Yuan Mark Liao. 2022. YOLOv7: Trainable bag-of-freebies sets new state-of-the-art for real-time object detectors. *arXiv preprint arXiv:2207.02696* (2022).
- [75] Patricia Warren, Donald Tomaskovic-Devey, William Smith, Matthew Zingraff, and Marcinda Mason. 2006. Driving while black: Bias processes and racial disparity in police stops. *Criminology* 44, 3 (2006), 709–738.
- [76] Benjamin Wilson, Judy Hoffman, and Jamie Morgenstern. 2019. Predictive inequity in object detection. *arXiv preprint arXiv:1902.11097* (2019).
- [77] Qian Xie, Dawei Li, Zhenghao Yu, Jun Zhou, and Jun Wang. 2020. Detecting Trees in Street Images via Deep Learning With Attention Module. *IEEE Transactions on Instrumentation and Measurement* 69, 8 (2020), 5395–5406. <https://doi.org/10.1109/TIM.2019.2958580>
- [78] Fisher Yu, Haofeng Chen, Xin Wang, Wenqi Xian, Yingying Chen, Fangchen Liu, Vashisht Madhavan, and Trevor Darrell. 2020. Bdd100k: A diverse driving dataset for heterogeneous multitask learning. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*. 2636–2645.
- [79] JD Zamfirescu-Pereira, Jerry Chen, Emily Wen, Allison Koenecke, Nikhil Garg, and Emma Pierson. 2022. Trucks Don't Mean Trump: Diagnosing Human Error in Image Analysis. In *2022 ACM Conference on Fairness, Accountability, and Transparency*. 799–813.
- [80] Hui Zhou, Han Wang, Handuo Zhang, and Karunasekera Hasith. 2020. LaCNet: Real-time end-to-end arbitrary-shaped lane and curb detection with instance segmentation network. In *2020 16th International Conference on Control, Automation, Robotics and Vision (ICARCV)*. IEEE, 184–189.