

“Data Strikes”: Evaluating the Effectiveness of a New Form of Collective Action Against Technology Companies

Nicholas Vincent
Northwestern University
Evanston, IL, USA,
nickvincent@u.northwestern.edu

Brent Hecht
Northwestern University
Evanston, IL, USA,
bhecht@northwestern.edu

Shilad Sen
Macalester College
Saint Paul, MN, USA,
ssen@macalester.edu

ABSTRACT

The public is increasingly concerned about the practices of large technology companies with regards to privacy and many other issues. To force changes in these practices, there have been growing calls for “data strikes.” These new types of collective action would seek to create leverage for the public by starving business-critical models (e.g. recommender systems, ranking algorithms) of much-needed training data. However, little is known about how data strikes would work, let alone how effective they would be. Focusing on the important commercial domain of recommender systems, we simulate data strikes under a wide variety of conditions and explore how they can augment traditional boycotts. Our results suggest that data strikes can be effective and that users have more power in their relationship with technology companies than they do with other companies. However, our results also highlight important trade-offs and challenges that must be considered by potential organizers.

CCS CONCEPTS

• **Human-centered computing~Human computer interaction (HCI)** • Information systems~Recommender systems

KEYWORDS

Data strikes, recommender systems, online collective action

ACM Reference format:

Nicholas Vincent, Brent Hecht and Shilad Sen. 2019. “Data Strikes”: Evaluating the Effectiveness of a New Form of Collective Action Against Technology Companies. In *Proceedings of WWW '19: The Web Conference (WWW '19)*, May 13, 2019, San Francisco, USA. ACM, New York, NY, USA, 13 pages. <https://doi.org/10.1145/3308558.3313742>

This paper is published under the Creative Commons Attribution 4.0 International (CC-BY 4.0) license. Authors reserve their rights to disseminate the work on their personal and corporate Web sites with the appropriate attribution.

WWW '19, May 13–17, 2019, San Francisco, USA
© 2019 IW3C2 (International World Wide Web Conference Committee), published under Creative Commons CC-BY 4.0 License.
ACM ISBN 978-1-4503-6674-8/19/05
<https://doi.org/10.1145/3308558.3313742>

1 Introduction

Large technology companies are facing a growing wave of public criticism. Just in the last year, these companies have been condemned for a wide range of practices, including those related to privacy [13, 55], harassment [20], addiction [10], effects on democracy [12], and automation [1]. The breadth and scale of the public concerns about tech companies has even led to the popularization of the term “Big Tech” [25], an adaptation of the terms “Big Oil” and “Big Tobacco” [3, 28].

However, these same companies that anger the public often are *dependent on the public* in new ways. Specifically, in addition to needing users and customers to generate revenue, *tech companies often rely on the public’s “data labor”* [44] to power *mission-critical intelligent technologies*. For example, Google requires user clicks to train its ranking algorithm [45]. Similarly, the highly-profitable recommender systems employed by companies like Amazon and Netflix require large amounts of data from users (i.e. ratings, clicks, and views) [17, 53].

Seen through the lens of the public’s concerns about tech companies, these companies’ dependence on user data to fuel their intelligent technologies can be understood as a potentially powerful source of new leverage for the public. To help the public action this leverage, several authors have proposed the notion of “**data strikes**” (e.g. [2, 39, 44, 52]), in which users halt their data labor [44]. The basic logic that motivates data strikes is straightforward: if users withhold their data labor from a tech company, some of the company’s essential services will suffer, and this would then force the company to make concessions that are desired by the public. These concessions could range from improved privacy policies to profit sharing [2, 15].

Despite the growing discussion around data strikes, little is known about how this new type of collective action would work or about how data strikes relate to standard forms of collective action like traditional consumer **boycotts** (a type of collective action classified as political consumption [32, 42]). Additionally, as data strikes increasingly enter the realm of feasibility (see below), there is little empirical information about how effective data strikes could be, let alone the data strike configurations that would be most effective. Activists seeking to organize a data strike have no guidance regarding the number of users that would need to join them, the kinds of services most vulnerable, the types of users that would allow them to be most successful, or even whether strikes can be successful at all. Similarly, tech

companies are not aware of the potential damage that could be inflicted through data strikes.

This paper seeks to improve our basic understanding of data strikes and provide much-needed empirical information about their effectiveness. We first situate data strikes in relationship to traditional boycotts, in which a user stops patronizing a company entirely. Through the introduction of a lightweight framework that (partially) describes **collective action** in a technology company context, we highlight that most traditional boycotts against a company operating data labor-dependent intelligent technologies will implicitly also include a data strike, but that data strikes can also occur independently from boycotts. For example, a consumer who continues to purchase products from an online retailer could engage in a data strike by using private browsing windows and not providing product ratings.

Next, focusing on the domain of recommender systems, a family of intelligent technologies that are critical drivers of revenue [17, 51], we introduce a novel evaluation procedure for understanding collective action campaigns against technology companies. Our procedure uses a metric called **surfaced hits**, which can capture the effects of both a traditional boycott and a data strike. Leveraging surfaced hits, we examine how model performance changes depending on 1) the size of the participating group, 2) whether the participating group is a random group of users or a homogeneous group of users who share some characteristics (e.g. women, people interested in documentary movies), and 3) whether or not the group is conducting an independent data strike or are doing so as part of a traditional boycott.

Our results confirm that users' **data labor power** – which is mostly unique to online platforms and is manifest in a data strike – provides users with a new source of leverage in their relationships with technology companies. For small recommenders and in specific product spaces, this added leverage can be particularly substantial. A moderately-sized data strike alone – even when not part of a traditional boycott – can significantly harm the performance of a recommender system. Indeed, for moderately-sized data strikes, we observe recommender accuracy decreasing to the levels that defined the state-of-the-art in recommender algorithms in 1999. This power comes from the reduced performance for both non-striking users (who receive recommendations trained on less data) and striking users themselves (who receive recommendations that are not personalized). Additionally, our work shows that data strikes that occur as part of a traditional boycott add data labor power to the standard **consumer power** from a boycott, increasing the overall power of the collective action campaign.

Finally, our work also highlights that data strikes that are not part of a traditional boycott represent a *fundamentally new type of collective action*, one in which the *barrier to entry is much lower than in a boycott*. Most notably, we observe that data strike participants can substantially reduce the utility of a recommender system without sacrificing access to the underlying products and services. Given that it has proven difficult for people with limited financial resources to participate in political consumption activities like boycotts [32], the

demonstrated effectiveness of data strikes could democratize access to these activities (e.g. users who cannot afford to use expensive alternatives to online platforms can still strike). This is analogous to an offline boycott in which a user who cannot afford expensive pizza could still participate in collective action against a local low-price pizza chain while continuing to buy their products.

Below, we adopt a standard structure to motivate, explain, and expand on our findings. We first cover related work, then discuss methods, followed by results. We close with a discussion of the issues identified in our results and by highlighting limitations.

2 Related Work

In this section, we describe how this research draws motivation from four areas in particular: the growing discussion related to data strikes, research on the relationships between tech companies and volunteer-created content, studies that generally seek to quantify the financial value of user data, and studies that look specifically at ways to manipulate recommender system outputs.

2.1 Data Strikes

This research was most directly motivated by growing calls for collective action campaigns that force changes in technology platforms by leveraging the value of user data to these platforms [2, 11, 16, 23, 37, 44]. These growing calls use different, potentially conflicting framings of *data as capital* or *data as labor* [1]. With regard to the former (data as capital), collective action is framed in terms of a boycott, in which users stop their consumptive activities (e.g. purchasing products through a web platform, using a social media platform) which in turn prevents the flow of their capital (data, related revenue like advertising revenue) to the platform. This framing is exemplified by very recent boycotts put into practice against Facebook (e.g. [42]). The data labor view suggests that data “unions” should protect the interests of those who produce data (i.e. users) [2, 16, 44]. Just as traditional labor unions have implemented (and threatened) strikes to gain leverage when collectively bargaining, Lanier and others [16, 23] have written that data unions might similarly engage in a “strike”. These authors point out that users can leverage their data in ways that resemble both traditional boycotts and strikes.

The diverse understandings of collective action campaigns that use data leverage needed to be integrated in order to make these campaigns concrete enough to simulate. Below in the Framework section, we enumerate one possible integration and use the corresponding framework to inform the design of our experiments.

Ideas about collective action campaigns that use data leverage often imagine a future in which people can “delete their data” from an online platform, and this future is becoming increasingly realistic thanks to developments and discussions in the policy domain. For instance, the European Union recently adopted the General Data Protection Regulation (GDPR) [30], which includes a provision ensuring the right to erasure. Barring

special circumstances (e.g. data critical to public health research), individuals covered by the GDPR will have the right to request that their personal data be deleted. As such, the GDPR potentially empowers activists to engage in more powerful data labor-related collective action than previously possible, in particular by erasing old data instead of just stopping the flow of new data. While it remains to be seen how often and how effectively the right to erasure will be used in practice – and how it might apply to campaigns that use data leverage specifically – the inclusion of this right highlights a large shift in regulatory practices towards data usage. The GDPR could trailblaze the way for similar or even stronger provisions by other regulatory bodies (e.g. California’s State Government [6]).

In keeping with this policy trajectory and with the typical vision of collective action campaigns that use data leverage, we simulate campaigns in which people can “delete their data” when they participate. However, our approach also applies to other contexts, for instance domains in which there is *no existing data* like reviews about a new television show or location data used to predict traffic (and less directly to contexts in which participants cannot delete data, but do not execute new data labor). Furthermore, very recent research suggests that simple tools like browser extensions may help web users successfully join web-based collective action campaigns with low overhead for the user, which could help stop the flow of implicit behavioral data [39].

Finally, it is important to note that the social science literature can give some guidance as to how large one could expect the campaigns we examine to be. Data from Europe and the U.S. found that between 28% and 35% of consumers had engaged in an act of political consumption [32, 42], which includes either boycotting or “buycotting” (aligning one’s purchases with a company that is perceived to align with one’s political preferences).

2.2 Tech Companies and Volunteer-Created Content

A related area of research that also helped motivate this study is work that has sought to understand the dependence of technology companies on volunteer-created content like Wikipedia articles. McMahon et al. [40] showed that Google search effectiveness drops substantially when Wikipedia links are silently removed from search results, highlighting how important Wikipedia is to the success of search engines. Outside of Google’s relationship to Wikipedia, Vincent et al. found that Stack Overflow and Reddit receive substantial benefits from Wikipedia in the form of impactful links and references [56]; they showed that these benefits come in the form of both increased engagement from users and advertising revenue. Although not originally intended as such, these studies can be seen as simulating a form of data-related campaign in which companies are somehow prevented from using Wikipedia content. Such campaigns would be highly unlikely given Wikipedia’s content license [58] and other factors, but the design of these studies helped to inform our methodological approach outlined below.

2.3 Financial Value of Data

This work was more generally motivated by a broad body of literature seeking to understand the financial value of data and highlighting the importance of this understanding. This research includes efforts to provide individual users with transparency into the value they create, such as the Facebook Data Valuation Tool created by González Cabañas et al. [18], as well as efforts to broadly understand the value of that data at a macroeconomic scale (e.g. the value of Wikipedia to GDP statistics [4]). On the policy front, the World Economic Forum has identified data as a new “asset class” and suggests that thinking about data economics demands a new understanding of the personal data ecosystem [49].

2.4 Recommender System Manipulation

Within the recommender system literature, there has been research into various ways recommender systems might be manipulated. For instance, prior work has examined how recommender systems might be “shilled”, i.e. misled so as to promote a particular product (e.g. [7, 36]). Like potential strikes, shilling attacks are an adversarial approach to manipulating the outputs of a recommender. As we explain below, our experiments specifically focus on campaigns that withhold data, so findings related to shilling are not directly applicable. However, in practice collective action participants might be able to adopt techniques from shilling, in which case this body of literature may be useful to both users and the companies against which they seek to gain leverage.

Recent research from Wen et al. explored recommender performance under conditions in which users filter some portion of their preference history to increase privacy and/or recommender accuracy [57]. Specifically, the authors found that users can filter out some of their preference history from the recommender while maintaining, or even improving, performance for implicit recommendations. This research has direct implications for data-related collective action: although users might be able to perform a “partial strike” by deleting some of their data without suffering decreased performance, these partial strikes are unlikely to be very effective, as in some cases, partial strikes by all users may improve population-level performance. In our experiments, we focus on directly simulating conditions in which users withhold *all* their data, and we also explicitly consider both data strikes and traditional boycotts. However, exploring the interplay between strikes, boycotts and data-filtering tools will be an important area of future research that is mutually beneficial to both problem contexts. In particular, the data filtering interfaces described by Wen et al. could be another outlet for users to actuate data strikes.

3 FRAMEWORK

As mentioned above, although collective action campaigns that use data leverage against technology companies have been discussed as a theoretical possibility, they have not been characterized in detail. Indeed, in the context of collective action

against technology companies, the distinction between data strikes, boycotts, and combinations of the two can be unclear.

In order to simulate data strikes, we need to first concretely define data strikes and their relationship to traditional boycotts. To do so, we turned to the divergent theoretical underpinnings of the boycott and strike terms. In a boycott, participants are *consumers* who cut off the flow of an asset (e.g. money from purchases) to a firm. In a strike, participants are *laborers* who stop performing work for a firm. Users of an online platform can therefore boycott the platform by refusing to use the platform as consumers (e.g. not buying from an e-commerce site, not visiting a news site or video site, etc.) and strike against the platform by refusing to provide data (e.g. deleting data, preventing the flow of new data by using private browsing features or other privacy techniques like ad blockers or Mozilla’s new Facebook Container [14]).

In most cases, *boycotts against tech companies implicitly include a data strike*. For instance, users of a video platform like YouTube who boycott (refuse to visit the website) are also implicitly conducting a data strike by cutting off the flow of behavioral data like views, likes, and comments. However, *it is often possible to participate in a data strike without boycotting the platform*. This occurs if someone continues to access a website but withholds data using privacy-preserving techniques (e.g. private browsing), leverages data management options made possible through data protection regulation, or – critically for our context – *refuses to comment on products, rate products, or review products*. The inverse, a boycott without a data strike, is less ecologically valid in the context of our work. Someone who boycotts Amazon products is unlikely to submit product reviews and ratings. While there are more nuanced situations in which this could occur, in this paper, we simulate boycotts in concert with a strike.

To put the nuanced relationship between boycotts and data strikes as defined above into better context, we consulted the literature to identify the various specific means by which these types of collective action campaigns could affect company revenue (e.g. [2, 44]). We identified four such pathways to revenue impacts:

- The **direct data labor effects** (e.g. algorithmic performance decreases leading to loss in sales)
- The **indirect data labor effects** (e.g. because algorithmic performance goes below some threshold, users quit the platform leading to loss in sales) [54].
- The **direct consumer effects** (e.g. people stop buying products or viewing ads) [32, 42].
- The **indirect consumer effects** (e.g. a large number of customers stop buying products or viewing ads, so there is a loss of economy of scale advantages) [50].

The **consumer effects** above (which make up **consumer power**) are those that exist in traditional boycotts and have been felt by targeted businesses since well before intelligent technologies came into common use. The **data labor effects** (which make up **data labor power**) are specific to collective action campaigns against companies that use data-hungry

intelligent technologies. While a traditional boycott only includes direct and indirect consumer effects, a simultaneous data strike and boycott includes all four of the above components.

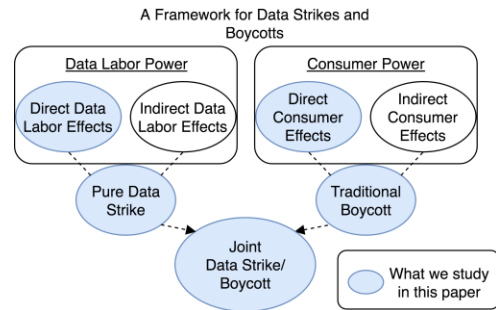


Figure 1: Graphical summary of our framework for defining data strikes and boycotts. Blue indicates aspects on which we focus.

This simple framework, depicted in Figure 1, provides a much-needed lens into the ontology of these types of collective action campaigns, but it also highlights that researchers – especially those outside a technology platform – can only simulate a portion of the effects of data strikes and boycotts. For instance, without exact sales numbers, both the direct and indirect sales effects are very difficult to study externally. Similarly, the indirect long-term effects of reduced recommender performance would be difficult to capture both for researchers external or internal to a platform (although there is at least one case of this being done internally in the search literature using an A/B test framework [54]).

Fortunately, in this paper, we show by focusing on the direct data labor effects while still considering the direct consumer effects, we can still learn a great deal about collective action campaigns that use data leverage. As we discuss below, our results support the effectiveness of data strikes, suggesting the unique relationship between users and technology companies can empower users beyond what would be the case in other contexts. We also discuss how our results can be interpreted as a *lower bound* for the effects of data strikes and boycotts against tech companies because we cannot measure indirect effects.

As discussed above, the data strikes we simulate correspond to a person “deleting their account” and using the recommender as a “guest”, with no account history. Similarly, the boycotts we simulate (which include a joint strike) correspond to someone deleting their account and not using the system at all. In both cases, users not participating in a strike or boycott receive recommendations from a model that has been trained without strikers’ or boycotters’ data. As we discuss in Limitations, there are many other configurations one can imagine – especially those related to strikes that do not involve the deletion of past data and play out over time – and we believe these to be important directions of future work.

4 METHODOLOGY

In this section, we first describe aspects of our methods that were consistent across all our experimental configurations. We then describe the two broad types of collective action campaigns we simulated: “general” groups comprised of randomly selected users and “homogenous groups” of users who share some characteristic (e.g. power users, fans of comedy movies).

4.1 Design of Experiments

In each experiment, we evaluate recommender systems under a variety of simulated data strike and boycott conditions. While the campaigning groups differ by experiment, the basic methods are the same.

4.1.1 Datasets. Our primary dataset was MovieLens-1M (ML-1M), which consists of 1 million “1 to 5 star” ratings for 3,706 movies provided by 6,040 users with self-reported demographic data [22]. To better understand performance against large recommenders, we additionally performed experiments with the much newer MovieLens-20M (ML-20M) dataset [22], which contains 20x more ratings but no demographic data about users. The MovieLens datasets have been hugely influential and have been central to recommender system research for decades [22].

4.1.2 Algorithm Choice and Implementation. For each experiment, we focus on the well-known and high-performing Singular Value Decomposition (SVD) recommender algorithm [41]. As validation, we also performed a smaller set of ML-1M experiments with an older and mathematically distinct algorithm: item-based K-Nearest Neighbors (k-NN) [48] adjusting for item and user baselines as described by Koren [33]. Both algorithms are implemented in the open-source Python library Surprise [27], which we extended for our experiments. All code used for our experiments and analyses is available for replication and extension on GitHub¹. We validated the accuracy of the SVD implementation by ensuring results were comparable to published results on the ML-1M dataset [26, 34, 38, 46]. These successful comparisons are summarized in the linked GitHub repository.

4.1.3 Evaluation Procedure. We evaluated the recommender with five-fold cross-validation; for each evaluation fold, 20% of total data is available for testing and each rating is tested in exactly one fold. While data that is held out because of a simulated campaign cannot be used for training, it can be used for testing. This means we can consider results from the perspective of striking users, who will receive non-personalized recommendations which are based on each movie’s average rating. This is a standard baseline for producing recommendations for users who lack personalized data and it is a widely-used way to provide preference information in a non-personalized fashion (e.g. displaying a movie’s average rating instead of predicted rating for a given user).

4.1.4 Metrics for Evaluating Strikes and Boycotts. When evaluating the accuracy of explicit rating predictions for recommender systems, one common approach is to measure the

error in individual predictions, e.g. through an accuracy metric such as root-mean-squared error (RMSE) which was used for the well-known Netflix Prize [41], or using information retrieval metrics such as gain (NDCG) [5] or precision. Retrieval metrics have been favored in recent years because of their ecological validity [8, 31] (e.g. “Top Ten Movies for You”).

However, these standard metrics only capture performance for users who receive recommendations, and do not capture the consumer effects of users leaving due to boycotts. Therefore, this approach is not well-suited to understand boycotts from the perspective of the system owner, because the loss in revenue from boycotting users will not be visible in traditional metrics. For instance, if we simulate an 80% boycott and measure the RMSE of predicted ratings for the remaining 20% of users, the change in RMSE does not account for the users who *left the system entirely*.

To understand the relationship between data strikes and boycotts, it is critical to capture both the direct consumer effects of boycotting users and the direct data labor effects of striking users. To do so, we introduce a new metric, which we call **surfaced hits**. The metric measures the fraction of hits (defined as a rating of at least 4.0, as is common in prior work, e.g. [34]) across an entire group of users (perhaps all users, or non-boycotting users). The underlying assumption, that one hit corresponds to one unit of value for a recommender system, is supported by the widespread use of analytic metrics such as click-through-rate in online systems [17]. A perfect algorithm will surface all hits, and therefore have a surfaced hits value of 1.0. This metric can be effectively viewed as a variant of precision that sums (rather than averages) across all users and sets individual thresholds for precision equal to how many positive ratings each user has. More explicitly,

for each user u :

t_u = u ’s test ratings

n_u = number of true ratings ≥ 4 in t_u

p_u = top n_u of t_u ordered by predicted rating

h_u = number of true ratings ≥ 4 in p_u

surfaced hits = $\text{sum}(\{h_u\}) / \text{sum}(\{n_u\})$

As stated earlier, for both data strikes and boycotts combined with data strikes, participating users’ ratings are withheld from all training data sets. Users who boycott contribute zero hits to the numerator of surfaced hits for their test ratings ($h_u = 0$). In other words, the surfaced hits value is “penalized” by marking all positive ratings for the user as a non-hit. For users in a data strike, we included the user’s test ratings in the calculation of surfaced hits, but “penalize” via non-personalized recommendations (i.e. movie averages) because the user’s training data is not available due to the strike.

Overall, the surfaced hits metric has three useful properties for studying data strikes and boycotts. First, when users boycott, surfaced hits is reduced proportionally to the number of positive ratings in the boycotting group. In other words, if enough users boycott to remove half of all the “hits” in the dataset, surfaced hits will be reduced by at least half. This allows us to understand the effects of strikes and boycotts from within a single reference frame. Second, this metric also accounts for differences in user

¹ https://github.com/nickmvincent/surprise_sandbox

behavior: a user with 1000 hits in their rating history has 10x the impact of a user with 100 hits in their history. This captures the potentially disproportional economic value of more active users. Third, we can calculate surfaced hits for different subsets of users (e.g. all users, non-striking users, users similar to striking users) to understand the effects of collective action from different perspectives. We verified that these metrics perform very similarly to well-established list metrics including NDCG and precision while at the same time capturing the damage that occurs when boycotting users leave a system (see Section 5.3 and the code repository).

4.2 Campaign Configurations

4.2.1 General campaigns. To get a general understanding of the relationship between campaign size and recommender system performance, we first simulated a series of “general” campaigns with random user selection. In these campaigns, the demographic make-up of the groups approximates the distribution of all users. We selected a sequence of 16 different group sizes ranging from 0.01% of users to 99% of users. For each group size, we randomly selected a group of that size to participate in the campaign. To reduce noise associated with different random configurations, we repeated each of these 16 experiments 250 times with a new random user sample for the ML-1M dataset and 40 times for the ML-20M dataset.

4.2.2 Homogeneous campaigns. We also simulated campaigns executed by “homogeneous” groups defined by shared patterns in rating behavior or shared demographic information. More specifically, we created five types of homogenous campaigns: campaigns by “fans” of specific movie genres, campaigns by three categories of demographic groups, and campaigns defined by rating behavior. We created groups of “fans” for each movie genre by identifying all users who rated at least ten movies of that genre and have an average rating for the genre of four or higher. To simulate demographically-defined campaigns, we created groups based on user-reported demographics, specifically male/female, age bracket, and occupation. For rating behavior campaigns, we created campaigns for “power users,” defined as the top 10% of raters and “low frequency” users, the bottom 10% of raters.

For each of the five types of homogeneous groups, we simulated campaigns in which 50% of users within a given group participated. For example, we simulated campaigns with groups such as 50% of all women, or 50% of all comedy movie fans, and so on. Importantly, 50% participation allows us to simulate what happens to similar users who do not participate. In other words, if some women participate in a data strike, what happens to women who do not? We also viewed 50% group participation as more realistic than full participation.

Our homogenous experiments focus on the data labor effects of data strikes. The consumer effects of a homogenous boycott scale with the size of the boycott as measured by the number of positive ratings, and as we show in our general campaign experiments, this effect is substantially larger than that of strikes, but this does not negate a strike’s value.

For each homogenous campaign configuration, we performed experiments with 50 sampled groups. We also compared the observed campaign effects to the expected effects for a random campaign with the same number of ratings. In order to obtain a relatively simple estimate of the “expected” effect of a data strike of some size, we computed a quadratic interpolation of our results shown in Figure 2. Our homogeneous experiments only consider the ML-1M dataset due to the lack of demographic information in ML-20M.

5 RESULTS

In this section, we first describe the relationship we observed between recommender performance and campaign size in the case of general (random users) collective action campaigns, focusing on comparing pure data strikes to strikes coupled with traditional boycotts and examining the overall effectiveness of campaigns across datasets. Next, we describe the key findings from our homogeneous group experiments, focusing on the finding that unique homogenous groups, defined demographically or behaviorally, can exert their data labor power to disproportionately affect similar users, indicating the potential for data strikes that target specific preference spaces to boost their effectiveness.

This section uses the surfaced hits metric, described above. We focus on the popular SVD algorithm because the item-based k-NN algorithm behaved similarly in our initial experiments (see GitHub repository).

5.1 General Campaign Experiments

We begin by examining the effect of general data strikes and boycotts (i.e. with random users) from the perspective of the system owner (e.g. Google, Facebook, operators of MovieLens). Next, still focusing on the perspective of system owners, we specifically zoom in on performance changes relative to un-personalized results.

Figure 2 shows the effect of data strikes (blue line) and joint data strikes and boycotts (green line) on surfaced hits (y-axis) across the system for both ML-1M (left column) and ML-20M (right column). As a reminder, a value of 1.0 would correspond to an algorithm that produces perfect ranked lists for every user.

These plots include dotted horizontal lines that provide important context: the black line shows performance of SVD with full access to the dataset (which gives 77.4% of hits), the red line shows the results of “MovieMean,” which gives completely un-personalized ratings (movies are ranked in order their mean rating) and the gold line shows the results of completely randomly ranked lists (i.e. worst-case performance). Note the high number of hits associated with random lists: this is because MovieLens users tend to give movies high ratings, so when evaluating even random lists many of the items suggested for a given user will be hits. We address this phenomenon below by focusing not only on raw changes in surfaced hits, but also on performance *relative to* un-personalized performance, i.e. we zoom in on performance change between the black and red lines.

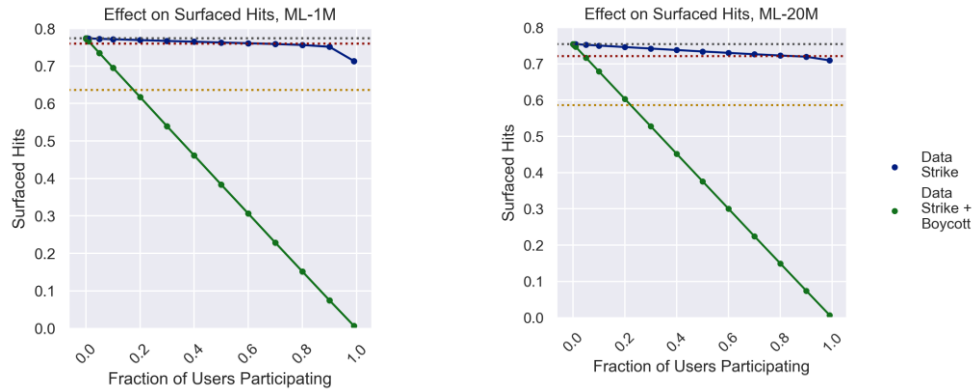


Figure 2: The relationship between campaign size and surfaced hits. Surfaced hits include all users (including strikers and boycotters), and therefore reflect the perspective of the system owner. Dotted horizontal lines provide comparisons: black (uppermost) shows fully personalized SVD, red (middle) shows un-personalized results, and gold (bottom) shows random results.

The most significant trend in Figure 2 is that boycotts are substantially more effective than data strikes. For instance, while a 30% boycott of ML-1M reduces hits to from 77.4% to 53.9%, a 30% data strike only reduces hits by 0.7% to 76.7% (ML-20M, right column, shows a similar trend). Furthermore, for both datasets a boycott of about 20% of users reduces surfaced hits to the amount expected for completely randomized recommendations. This result means that at first glance, the loss in hits caused by users who leave the system strongly outweighs the loss in hits from reduced algorithmic performance. Importantly, this finding does not fully explicate the potential of data strikes, as we will describe below.

We note that the gaps in Figure 2 between un-personalized results (red lines) and fully personalized results (black lines) appear to be small and correspond to a loss of 1.4% of surfaced hits for ML-1M and 3.3% for ML-20M. This reflects the non-linear value of recommender algorithms; the small margin between un-personalized and personalized algorithms corresponds to a large

amount of value for a recommender system operator. For example, in Netflix’s case, the margin between un-personalized algorithms and personalized algorithms accounts for a 2-4x increase in engagement with recommended items and \$1 billion in revenue [17] (see below). Thus, we now specifically focus on the change in performance relative to un-personalized results to inspect how data strikes leverage direct data labor effects to lower recommendation performance towards un-personalized levels.

Figure 3 zooms in on surfaced hits during a data strike using the same y-axis as Figure 2. Again, the black horizontal line marks personalized performance and the red horizontal line marks un-personalized performance. Additionally, in Figure 3 the horizontal cyan line shows the performance of simple item-based k-NN (which we evaluated with full access to each dataset), a technique that was introduced in 1999. This context shows the ability of campaigns to essentially set recommender system performance “back in time”.

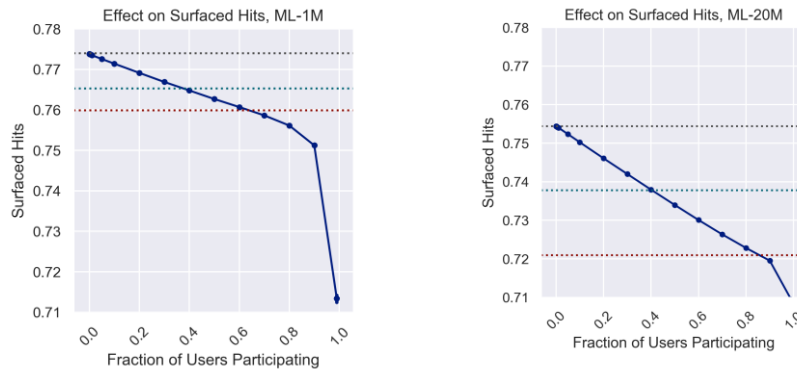


Figure 3: The relationship between data strike size and surfaced hits. Dotted horizontal lines provide comparisons: black (top) shows fully personalized SVD, cyan (middle) shows item-based k-NN (1999), and red (bottom) shows un-personalized results.

In Figure 3, the potential power of data strikes becomes clearer. The left side of the figure shows that campaigns had substantial effects on recommender performance for ML-1M relative to non-personalized ratings. For instance, a strike with 30% of users (which is a realistic size based on research on political consumption; see Related Work) degrades performance such that users lose roughly half the benefits of personalization. These results also illustrate the power of collective action to potentially negate decades of algorithmic advances. Looking again at MovieLens-1M, a strike by 37.5% of users can roll back hits to a level equivalent to the classic item-based k-NN algorithm introduced in 1999 [48] (cyan dotted line in Figure 3).

The ML-20M results (right column), however, suggest a somewhat more complicated story for recommenders using larger datasets. When the dataset size is increased by a factor of 20, a strike by the same percentage of users becomes somewhat less effective. If 30% of ML-1M strike (1800 users), we see a 50.2% reduction in the benefits of personalization, but if 30% of ML-20M strike (41400 users), it would only cause a 37.0% reduction.

A likely explanation of why the relationship between strike size and strike power differs between ML-1M and ML-20M lies in the two effects of a data strike, the effect on the strikers themselves and the effect on non-strikers. The first effect captures how the removal of the strikers' data lowers performance of the recommender for non-strikers (i.e. the ability of strikers to affect the experience of non-strikers). While we see similar directional relationships between ML-1M and ML-20M, this effect is more pronounced for ML-1M. For instance, at 30% strike participation, we see a 25.4% reduction in personalization for non-striking ML-1M users but just a 4.2% reduction in personalization for ML-20M (analysis in GitHub repo). ML-20M does not see an equivalent reduction in personalization for non-striking users until strike participation rates hit 77%. The explanation for this difference is likely straightforward: ML-20M has more redundant encodings of preference patterns due to its sheer size, so an equivalent percentage of strikers cannot have as much of an effect on the experience of non-striking users.

The second factor driving the strike results is the fact that, during a data strike (rather than boycott), by definition, striking users can still use the system. Because *these users must still receive a ranked list of items* (e.g. to power a Netflix-style interface), what used to be a personalized ranked list must now become a non-personalized ranked list. As noted above (Section 4.1.3), we implement this non-personalized ranked list in the most ecologically valid way possible: using the average movie rating from other users (item mean). As such, for each new user who strikes, the recommender will inch towards item mean by default, regardless of the effect on non-striking users. In other words, *even in the face of large amounts of training data, users can hurt the system by refusing to provide the input data needed to make predictions for themselves.*

Finally, we reiterate that evidence from industry suggests that in commercial systems, a change in surfaced hits has a non-linear value to recommender operators. In other words, the visually-small performance change between the red and black horizontal lines in Figures 2 and 3 may have an outsized effect

on platform revenue. As mentioned before, the small surfaced hits improvement due to personalization may correspond to a 2-4x increase engagement with recommended items in other contexts [17]. Similarly, in 2010 YouTube published findings that suggest recommendations add substantial value: almost 30% of video views came from their recommender system and the recommender was the main source of views for most videos [59]. An industry report from consulting company McKinsey estimates that recommender systems account for 35% of Amazon purchases and 70% of Netflix views [24]. Taken together, this means that in many contexts, the revenue effects of data strikes seen above would be magnified relative to their appearance in the figures.

5.2 Homogeneous Campaign Group Experiments

In our homogenous campaign experiments, we examine what would occur if a demographics- or taste-defined group engaged in a data strike (e.g. women or documentary fans). Specifically, for the reasons defined above, we simulate the effects of 50% of the group striking, examining the impact on the remaining 50% in the group, as well as on the recommender overall. We use the term **Similar Users** to describe non-striking members of the striking campaign group (e.g. non-striking women in a strike by women) and **Not Similar Users** for all other non-striking users (e.g. non-women in a strike by women). Then, we define the **Similar User Effect Ratio** as the percent change in surfaced hits for Similar Users divided by the percent change in surfaced hits for Not Similar Users. One challenge with analyzing these homogenous boycotts is that the various groups are very different in size (see Table 1). Therefore, this set of analyses focuses specifically on percent change in surfaced hits, which partially mitigates the challenge in comparing data strikes by groups that are very different in size. Furthermore, when looking at the aggregate effects of homogeneous campaigns, we specifically look at the perspective of non-participating users, as both the direct consumer effects and the effects from striking users seeing un-personalized recommendations are functions of campaign size.

The results from our homogenous campaign experiments show that data strikes may be especially effective at impacting recommendations within targeted topical domains. Specifically, we observe that if homogeneous groups of people strike, non-striking users that share the same characteristic as the homogenous striking group will experience larger reductions in recommender accuracy than other users. For example, striking horror movie fans can make a large movie recommender suffer for other horror fans, potentially giving a competing movie site an opening. A secondary, related observation from these experiments, which we return to at the end of this section, is that homogenous groups' aggregate effects on non-participating users are not consistent: some groups have an outsized aggregate effect relative to their size, and vice versa.

Table 1: Examples of homogenous groups, the number of ratings in each group, percent change in surfaced hits, and the Similar User Effect Ratio.

Name	# Ratings	% change surfaced hits, Similar Users	% change surfaced hits, Not Similar Users	Similar User Effect Ratio
men	753769	-0.71	-0.64	1.11
women	246440	-0.38	-0.09	4.24
fans of horror	48464	-0.24	-0.03	7.16
under 18	27211	-0.27	-0.03	9.02
25-34	395556	-0.38	-0.3	1.28
56+	38780	-0.18	-0.03	6.97
artist	50068	0.15	-0.07	-2.26
power users	381407	-0.76	-0.45	1.71

Table 1, which includes a variety of example groups, highlights the primary finding that some groups are especially effective at lowering performance for Similar Users compared to other users. For instance, looking at Table 1, we see that when half of women strike, surfaced hits decreased for non-striking women by 0.38% while surfaced hits decreased for non-women by 0.09%. Therefore, the Similar User Effect Ratio is $0.38 / 0.09 = 4.24$. This effect is exaggerated even further in the case of the “under 18” group, which has a Similar User Effect Ratio of 9.02. On the other hand, the “25-34” age group has a more or less “flat” ratio of 1.28, and the same is the case for men.

It is clear from Table 1 that some homogenous groups are able to “punch above their weight” with respect to affecting the experience of non-striking users, at least those within their homogenous group. This is particularly important with respect to large recommenders, which we saw above are more robust to this component of the data strike effect (the effect of strikers on the experience of non-strikers). Regardless of the scale of data resources available to a recommender, it appears that recommendation quality for some users may always be vulnerable to targeted data strikes by Similar Users.

One hypothesis that explains why some data strikes hurt Similar Users more than other users lies in a holistic view of the user preference space. If a group of users has substantial uniqueness— or more specifically, mathematical independence – in their preferences when compared to other groups, a campaign by that group is less likely to hurt users not in that group. At the extreme, a group whose preferences are completely orthogonal to every other group may be able to execute a campaign without substantially affecting personalization for any other group (one realistic example might be groups based on language proficiency).

To understand this relationship, we compared a group’s Similar User Effect Ratios from Table 1 to a measure of preference independence based on overlap in rate/no-rate behaviors. To calculate preference independence, we first create

a vector with a column for each movie (3706 columns) and value equal to the proportion of users in the group who have rated the movie (i.e. the group implicit rating vector). A group’s preference independence is the cosine distance of that group’s vector to the similarly calculated vector for the all groups (i.e. the centroid). We focus on groups with over 20k ratings, ignoring the six very small groups for which Similar User effects are extremely noisy.

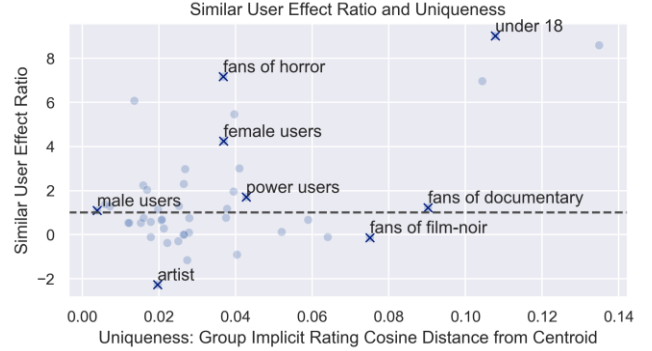


Figure 4: Scatterplot showing how homogeneous data strikes (with > 20k ratings) affect similar users differently than the general population. Along the x-axis, groups are organized by increasing uniqueness, defined by the cosine distance between the group implicit rating vector for the group and the general population. The y-axis shows the Similar User Effect Ratio. Gray dotted line shows ratio of 1. Pearson correlation is 0.55. “X” markers indicate labeled examples.

Figure 4 shows a full scatterplot of Similar User Effect Ratio for all groups with over 20k ratings. On the x-axis, we plot preference independence as defined above and on the y-axis, we plot Similar User Effect Ratio (note that negative values correspond to strikes that improve surfaced hits for Similar Users). Our experiments suggest a moderate positive relationship between a group’s preference independence and the effect of a strike on non-participating Similar Users (Pearson’s $r = 0.55$, $p < 0.001$), although this is likely driven by a non-trivial number of strong outliers (e.g. “under 18”, “fans of horror” in Figure 4) given the Spearman correlation ($r = 0.18$, $p > 0.05$) is not statistically significant. The full interplay between preference independence and strike effectiveness is a fertile ground for research, and future work could more closely study how preference spaces might be operationalized for the purpose of data strikes.

The outsized Similar User effects for many homogenous strikes point to both strategic advantages and effects that may limit adoption. For example, if women data strike against a company to affect some lasting change (e.g. ending discriminatory practices, launching profit sharing), a homogenous strike that focuses on recruiting other women may be especially effective, with large Similar User strike effects potentially driving women away from the company to competitors and substantially reducing company revenue. However, these effects may decrease strike adoption in some

contexts. For instance, returning to a strike by women against a major tech company, imagine that the tech company provides important services to large populations (e.g. online shopping in areas with few brick-and-mortar stores, low-cost communication in areas without similarly priced options). Many potential participants may be unwilling to participate in a campaign that will disproportionately damage these services for other women, thus limiting the adoption of homogenous campaigns in certain contexts.

A related secondary observation from these experiments is that homogenous groups have varying effects on all non-strikers (rather than just Similar Users) compared to the “expected” effects based on their number of ratings (determined by a quadratic interpolation of our general results, as described above). In other words, homogenous strikes vary in their ability to “punch above their weight” with respect to their capacity to decrease performance for all non-participating users: In 34 of 50 homogeneous groups, surfaced hits for non-participating users decreased more than would be expected for a random group with an equal number of ratings. Ratios of observed to expected effect ranged from 0.52 for the group of users whose occupation was “retired” to 1.84 for “power users”, with other examples including a 0.7 ratio for “fans of comedy” and a 1.35 ratio for “fans of fantasy” (see GitHub repository for more details).

Based on this secondary finding, it seems that homogenous campaigns will not always be effective at damaging the general population of users because some groups are under-performing, and even over-performing groups are limited by size in their ability to affect the general population. However, organizing campaigns with preference spaces in mind (i.e. campaigns that are homogenous in topical preferences) likely will be effective and, critically, may provide a way to challenge recommenders when the number of strikers is not sufficient to cause more general damage to performance.

5.3 Generalizing Beyond “Surfaced Hits” and SVD

In presenting our results, we have focused on our “surfaced hits” metrics. However, we also computed the more typical recommender systems evaluation metrics of RMSE, NDCG with all items, NDCG@k, Precision@k, and Recall@k for $k = \{5, 10\}$. We additionally calculated these metrics when only including “long-tail” (i.e. unpopular) movies. All these metrics produce similar results regarding the effects of various data strikes configurations, although they did not afford us the ability to analyze boycotts. The full dataset of results is available in our GitHub repository. We also note that in our early experiments using item-based k-NN, ML-1M, and traditional metrics, we observed very similar general trends, e.g. the effect of data strikes on the NDCG@10 of an item-based k-NN recommender mirrored effects on other traditional metrics for SVD.

6 DISCUSSION

In this paper, we have taken the previously hypothetical notion of data strikes, identified a wide variety of realistic campaign

configurations (including when they are combined with traditional boycotts), and simulated the effect of these configurations in the recommender systems domain using best-practice evaluations. Comparing these campaign configurations and looking specifically at direct effects, we find that while the consumer power of boycotts still substantially outweighs the data labor power of data strikes, data strikes represent a potentially powerful new form of leverage. Moreover, we saw that strike organizers might specifically target preference spaces within a recommender to achieve especially effective data strikes within those spaces. Below, we discuss some of the more general implications of these results.

6.1 Barriers to Entry vs. Impact

Our results suggest that collective action organizers targeting companies that operate intelligent technologies like recommender systems have more options than is the case in traditional collective action. Specifically, these organizers have the ability to optimize for impact or for barrier to entry. Our results show that boycotts have a larger impact, but they require all participants bear the cost of not using a potentially valuable service. Strikes, on the other hand, allow users to continue to benefit from the use of technology platforms without completely losing the ability to collectively bargain. Historically, participation in political consumption has been easier for affluent groups [32] - data strikes represent an approach that may be substantially more accessible. Notably, our results suggest that new technologies focused on privacy (e.g. initiatives from Mozilla [14]) and online political consumption (e.g. recent work from Li et al. [39]) are a promising approach to empowering more individuals to participate in collective action.

This notion of low-barrier-to-entry collective action echoes early research on digital activism by Earl and Kimport, who argued that the web reduces costs for participating in protest behavior like boycotts, petitions, and email campaigns [9]. Specifically, data strikes can be seen as another tool in the toolbox of low-barrier-to-entry techniques, offering an additional avenue for people to take action against technology companies. However, we also observed that this increased accessibility is coupled with reduced power.

6.2 Towards a Holistic View of Data Strikes

Although our results give us important insight into the potential impact of data strikes and boycotts, our work likely only captures a portion of the real-world effect of a collective action against an intelligent technology. In particular, as is discussed in the Framework section, we cannot measure directly the indirect effects of traditional boycotts or data strikes. This means that our results should be interpreted as a lower bound on the effects of any data labor-related collective action campaign.

A related point that emerges from our results viewed with the lens of our framework is that collective action against technology companies will largely be more powerful than collective action against non-technology companies. The effect of a boycott against, for instance, a clothing company, would largely not include either direct or indirect loss of data labor

value (excluding edge cases like long-term sales and marketing data). Since our results suggest that these factors will be non-trivial in most tech company boycotts, a user boycotting a tech company is likely to have a greater effect on revenue than would be expected in a boycott with a more traditional type of target. We expand on these power dynamics further below.

6.3 The Power of Algorithms vs. The Power of Public Data

Outside of the context of boycotts and data strikes, our results can also be viewed as a means to better understand the power of data provided by the public relative to the power of algorithms. Namely, we observed that moderately-sized strikes can bring recommender accuracy down to the levels of early recommender systems from 1999. These results, along with the work of McMahon et al. [40] and others [18, 23], emphasize the data leverage that the public has in its relationship with data-hungry intelligent technologies and the companies that operate them. While the public perception of intelligent technologies like recommender systems is that they are largely the accomplishment of tech companies and the computer scientists they employ, these technologies are in fact a highly cooperative project between the public and companies. Without the companies and computer scientists, the intelligent technologies do not exist. But the same is also true for the public's contributions of data (i.e. data labor). This implies a much different power dynamic than is currently assumed by most people on both sides of this relationship.

6.4 Limitations

This work has several limitations not yet discussed above that should be highlighted. First and foremost, this paper focused on recommender systems which, while a business-critical family of intelligent technologies, is only one family of intelligent technologies that could be vulnerable to collective action campaigns. Future work should seek to replicate our research for other intelligent technologies, for instance search ranking algorithms (e.g. [45]), "newsfeed"-style technologies (e.g. [43]), traffic prediction (e.g. [21]), and wi-fi geolocation (e.g. [19]).

As noted above, we simulate boycotts in concert with strikes owing to that being the more ecologically valid choice in the context of our study. We note that one could imagine boycotts coupled only with partial strikes: e.g. someone who boycotts a system but does not delete their past ratings. Exploring these types of configurations - and longitudinal considerations in general - is an important direction of future work.

While we used best-practice evaluation techniques in the recommender systems community [8, 22, 46], these techniques have several limitations that also affect the large literature of recommender systems research that employs them. In particular, we considered only explicit ratings and did not consider implicit preferences expressed through user behavior (which are not available in the MovieLens dataset). We also only considered our recommenders in an offline environment (as opposed to in a live experiment). Finally, to gain more insight into the nuances of

recommenders, it will be valuable to explore other recommender system datasets, particularly datasets from industry contexts.

Another important limitation is that in our experiments, we had to operationalize male/female as a binary variable due to the data available in the MovieLens dataset. Similarly, we were not able to test other types of demographic groups (e.g. LGBT communities, political groups). Relatedly, our use of the term "homogenous" refers to a specific demographic or topical dimension; it does not consider the diversity within our "homogenous" groups, and doing so would be a fruitful area of future work.

Finally, this paper focused on understanding the effect of collective action campaigns of various sizes and types, but it did not consider the collective action problem of organizing or actuating these campaigns. Fortunately, this problem maps to a deep body of work within social computing and related fields on sociotechnical strategies for motivating collective action online (e.g. [35, 47]). An obvious direction of future work in this research space involves building tools to organize data strikes and boycotts that leverages this body of work (either using GDPR or restricting new data collection). Recent research suggests that user-friendly tools like browser extensions may be an effective approach for making collective action campaigns easy to join and conduct [39].

6.5 Potential Negative Impacts

In response to calls for the computing community to better engage with the negative impacts of our research [29], we wish to highlight two major concerns with this work. First, we emphasize that our findings may be equally useful to organizers of campaigns as to they are to companies interested in mitigating the effectiveness of such campaigns. Relatedly, it is entirely possible that using a simulated data strike methodology, companies could identify which groups of users are and are not "useful to the algorithm", i.e. they could rank groups based on their utility to some intelligent technology and use this ranking to justify ignoring the interests of some groups. If this occurred along demographic lines, this could lead to troubling societal outcomes, e.g. if majority groups can collectively bargain with tech companies and minority groups cannot. Technologies to organize data strikes and boycotts could help mitigate this issue by recruiting users from a variety of demographic groups (perhaps, guided by future work, specifically targeting some preference space like Comedy movies or electronics products). Our results suggest that this should be a priority in the design of these technologies.

Moreover, our ability to perform simulated campaigns was predicated on the public availability of the MovieLens dataset. Substantially more accurate simulations could be run using much richer datasets available only to corporations, so in any "data strike simulation arms race", there will be a clear advantage for corporations. This means that corporations may be able to prepare models in advance to counteract boycotts or strikes. This might be mitigated through crowdsourced data collection or other means, a ripe area for future work.

7 CONCLUSION

In this paper, we have done the work of advancing the notion of data strikes from abstract discussion point to concrete campaigns that can be simulated. Through these simulations, we provided critical early empirical information to help advance the discussion around data strikes. In doing so, we first detailed a framework that describes the key elements of collective action campaigns that action data leverage (i.e. data strikes and traditional boycotts, which respectively leverage *data labor power* and *consumer power*). We found that these campaigns can be effective, with relatively small strikes wiping away significant portions of the value of recommender systems relative to simpler techniques. However, as datasets grow larger, data strikes become less effective, and strategies that target specific groups of users or preference spaces may become necessary. We discussed the implications of our results for those seeking to organize data strikes and companies seeking to understand potential effects on their core functionality.

8 ACKNOWLEDGMENTS

We are thankful to our anonymous reviewers, Loren Terveen, Rohan Pitre, Michael Ekstrand, and Lisa Lendway for the comments and feedback on the work.

REFERENCES

- [1] 1 million jobs will disappear by 2026. How to prepare for automation-commentary: 2018. <https://www.cnn.com/2018/02/02/automation-will-kill-1-million-jobs-by-2026-what-we-need-to-do-commentary.html>.
- [2] Arrieta Ibarra, I. et al. 2018. Should We Treat Data as Labor? Moving Beyond "Free." *American Economic Association Papers & Proceedings*. 1, 1 (2018).
- [3] Big Tobacco - Wikipedia: https://en.wikipedia.org/wiki/Big_Tobacco.
- [4] Brynjolfsson, E. and McAfee, A. 2014. *The second machine age: Work, progress, and prosperity in a time of brilliant technologies*. WW Norton & Company.
- [5] Burges, C. et al. 2005. Learning to Rank Using Gradient Descent. *Proceedings of the 22Nd International Conference on Machine Learning* (New York, NY, USA, 2005), 89–96.
- [6] California just passed one of the toughest data privacy laws in the country - The Verge: 2018. <https://www.theverge.com/2018/6/28/17509720/california-consumer-privacy-act-legislation-law-vote>.
- [7] Chirita, P.-A. et al. 2005. Preventing shilling attacks in online recommender systems. *Proceedings of the 7th annual ACM international workshop on Web information and data management* (2005), 67–74.
- [8] Cremonesi, P. et al. 2010. Performance of Recommender Algorithms on Top-n Recommendation Tasks. *Proceedings of the Fourth ACM Conference on Recommender Systems* (New York, NY, USA, 2010), 39–46.
- [9] Earl, J. and Kimport, K. 2011. *Digitally enabled social change: Activism in the internet age*. MIT Press.
- [10] Early Facebook and Google Employees Form Coalition to Fight What They Built - The New York Times: 2018. <https://www.nytimes.com/2018/02/04/technology/early-facebook-google-employees-fight-tech.html>.
- [11] Facebook campaign urges users to boycott Facebook for a day: 2018. <https://www.theguardian.com/technology/2018/apr/07/facebook-campaign-urges-users-boycott-facebook-for-one-day-protest-cambridge-analytica-scandal>.
- [12] Facebook admits social media sometimes harms democracy - The Washington Post: 2018. https://www.washingtonpost.com/news/the-switch/wp/2018/01/22/facebook-admits-it-sometimes-harms-democracy/?utm_term=.f0cf046c1ebe.
- [13] Facebook and Cambridge Analytica: What You Need to Know as Fallout Widens - The New York Times: 2018. <https://www.nytimes.com/2018/03/19/technology/facebook-cambridge-analytica-explained.html?mtrref=www.google.com&gwh=DB24EEF2DB7C9762BA867E70C9C6DB2C&gwt=pay>.
- [14] Facebook Container Extension: Take control of how you're being tracked | The Firefox Frontier: 2018. <https://blog.mozilla.org/firefox/facebook-container-extension/>.
- [15] Granholm, J. and Eldred, C. Facebook owes you money (opinion) - CNN: 2018. <https://www.cnn.com/2018/04/11/opinions/facebook-should-pay-us-for-using-our-data-granholm-eldred/index.html>.
- [16] Facebook users unite! "Data Labour Union" launches in Netherlands | Reuters: 2018. <https://www.reuters.com/article/us-netherlands-tech-data-labour-union/facebook-users-unite-data-labour-union-launches-in-netherlands-idUSKCN1IO2M3>.
- [17] Gomez-Urbe, C.A. and Hunt, N. 2015. The Netflix Recommender System: Algorithms, Business Value, and Innovation. *ACM Trans. Manage. Inf. Syst.* 6, 4 (Dec. 2015), 13:1–13:19. DOI: <https://doi.org/10.1145/2843948>.
- [18] González Cabañas, J. et al. 2017. FDVT: Data Valuation Tool for Facebook Users. *Proceedings of the 2017 CHI Conference on Human Factors in Computing Systems* (2017), 3799–3809.
- [19] Google, Apple tap crowdsourcing to map out WiFi locations | IT Business: 2010. <https://www.itbusiness.ca/news/google-apple-tap-crowdsourcing-to-map-out-wifi-locations/15658>.
- [20] Google sees major claims of harassment and discrimination as lawsuits proceed | Technology | The Guardian: 2018. <https://www.theguardian.com/technology/2018/mar/28/google-sexual-harassment-pay-gap-lawsuits-proceed>.
- [21] Graunke, G.L. 2001. *Using predictive traffic modeling*. Google Patents.
- [22] Harper, F.M. and Konstan, J.A. 2015. The MovieLens Datasets: History and Context. *ACM Trans. Interact. Intell. Syst.* 5, 4 (Dec. 2015), 19:1–19:19. DOI: <https://doi.org/10.1145/2827872>.
- [23] Hecht, B. 2017. HCI and the U.S. Presidential Election: A Few Thoughts on a Research Agenda. *CHI '18 Panel Presentation* (Denver, CO, 2017).
- [24] How retailers can keep up with consumers: 2013. <https://www.mckinsey.com/industries/retail/our-insights/how-retailers-can-keep-up-with-consumers>.
- [25] How Silicon Valley became "Big Tech.": http://www.slate.com/articles/technology/technology/2017/11/how_silicon_valley_became_big_tech.html.
- [26] Huang, S. et al. 2015. Listwise Collaborative Filtering. *Proceedings of the 38th International ACM SIGIR Conference on Research and Development in Information Retrieval* (New York, NY, USA, 2015), 343–352.
- [27] Hug, N. 2017. *Surprise, a Python library for recommender systems*. <http://surpriselib.com>.
- [28] Inside the Big Oil Game - TIME: <http://content.time.com/time/magazine/article/0,9171,920328,00.html>.
- [29] It's Time to Do Something: Mitigating the Negative Impacts of Computing Through a Change to the Peer Review Process - ACM FCA: <https://acm-fca.org/2018/03/29/negativeimpacts/>.
- [30] Key Changes with the General Data Protection Regulation: <https://www.eugdpr.org/the-regulation.html>.
- [31] Koenigstein, N. 2017. Rethinking Collaborative Filtering: A Practical Perspective on State-of-the-art Research Based on Real World Insights. *Proceedings of the Eleventh ACM Conference on Recommender Systems* (2017), 336–337.
- [32] Koos, S. 2012. What drives political consumption in Europe? A multi-level analysis on individual characteristics, opportunity structures and globalization. *Acta Sociologica*. 55, 1 (2012), 37–57.
- [33] Koren, Y. 2010. Factor in the Neighbors: Scalable and Accurate Collaborative Filtering. *ACM Trans. Knowl. Discov. Data*. 4, 1 (Jan. 2010), 1:1–1:24. DOI: <https://doi.org/10.1145/1644873.1644874>.
- [34] Koyejo, O. et al. 2013. Retargeted Matrix Factorization for Collaborative Filtering. *Proceedings of the 7th ACM Conference on Recommender Systems* (New York, NY, USA, 2013), 49–56.
- [35] Kraut, R.E. et al. 2012. *Building successful online communities: Evidence-based social design*. MIT Press.
- [36] Lam, S.K. and Riedl, J. 2004. Shilling recommender systems for fun and profit. *Proceedings of the 13th international conference on World Wide Web* (2004), 393–402.
- [37] Lanier, J. 2014. *Who owns the future?*. Simon and Schuster.
- [38] Lawrence, N.D. and Urtasun, R. 2009. Non-linear Matrix Factorization with Gaussian Processes. *Proceedings of the 26th Annual International Conference on Machine Learning* (New York, NY, USA, 2009), 601–608.
- [39] Li, H. et al. Out of Site: Empowering a New Approach to Online Boycotts. *Proceedings of the 2018 Computer-Supported Cooperative Work and Social Computing (CSCW'2018 / PACM)*.
- [40] McMahon, C. et al. 2017. The Substantial Interdependence of Wikipedia and Google: A Case Study on the Relationship Between Peer Production Communities and Information Technologies. *ICWSM* (2017), 142–151.
- [41] Netflix Update: Try This at Home: <http://sifter.org/simon/journal/20061211.html>.
- [42] Newman, B.J. and Bartels, B.L. 2011. Politics at the checkout line: Explaining political consumerism in the United States. *Political Research Quarterly*. 64, 4 (2011), 803–817.

- [43] Paek, T. et al. 2010. Predicting the Importance of Newsfeed Posts and Social Network Friends. *Proceedings of the Twenty-Fourth AAAI Conference on Artificial Intelligence* (Atlanta, Georgia, 2010), 1419–1424.
- [44] Posner, E.A. and Weyl, E.G. 2018. *Radical Markets: Uprooting Capitalism and Democracy for a Just Society*. Princeton University Press.
- [45] Radlinks, F. and Joachims, T. 2005. Query Chains: Learning to Rank from Implicit Feedback. *KDD '05: 11th ACM Conference on Knowledge Discovery and Data Mining* (Chicago, IL, 2005).
- [46] Said, A. and Bellogin, A. 2014. Comparative Recommender System Evaluation: Benchmarking Recommendation Frameworks. *Proceedings of the 8th ACM Conference on Recommender Systems* (New York, NY, USA, 2014), 129–136.
- [47] Salehi, N. et al. 2015. We Are Dynamo: Overcoming Stalling and Friction in Collective Action for Crowd Workers. *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems* (New York, NY, USA, 2015), 1621–1630.
- [48] Sarwar, B. et al. 2001. Item-based Collaborative Filtering Recommendation Algorithms. *Proceedings of the 10th International Conference on World Wide Web* (New York, NY, USA, 2001), 285–295.
- [49] Schwab, K. et al. 2011. Personal data: The emergence of a new asset class. *An Initiative of the World Economic Forum* (2011).
- [50] Shapiro, C. et al. 1999. Information rules: a strategic guide to the network economy. *Journal of Economic Education*. 30, (1999), 189–190.
- [51] Sharma, A. et al. 2015. Estimating the Causal Impact of Recommendation Systems from Observational Data. *Proceedings of the Sixteenth ACM Conference on Economics and Computation* (New York, NY, USA, 2015), 453–470.
- [52] Should internet firms pay for the data users currently give away? - The digital proletariat: 2018. <https://www.economist.com/news/finance-and-economics/21734390-and-new-paper-proposes-should-data-providers-unionise-should-internet>.
- [53] Smith, B. and Linden, G. 2017. Two decades of recommender systems at Amazon. com. *IEEE Internet Computing*. 21, 3 (2017), 12–18.
- [54] Song, Y. et al. 2013. Evaluating and Predicting User Engagement Change with Degraded Search Relevance. *Proceedings of the 22Nd International Conference on World Wide Web* (New York, NY, USA, 2013), 1213–1224.
- [55] The Latest: Facebook users await word on privacy scandal - The Washington Post: 2018. https://www.washingtonpost.com/national/the-latest-facebook-users-await-word-on-privacy-scandal/2018/04/09/37e8d82a-3c2a-11e8-955b-7d2e19b79966_story.html?utm_term=.1247b93f6998.
- [56] Vincent, N. et al. 2018. Examining Wikipedia With a Broader Lens: Quantifying the Value of Wikipedia's Relationships with Other Large-Scale Online Communities. *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems* (2018), 566.
- [57] Wen, H. et al. 2018. Exploring Recommendations Under User-Controlled Data Filtering. *ACM Conference on Recommender Systems (RecSys' 18)* (2018).
- [58] Wikipedia: Reusing Wikipedia content - Wikipedia: https://en.wikipedia.org/wiki/Wikipedia:Reusing_Wikipedia_content.
- [59] Zhou, R. et al. 2010. The impact of YouTube recommendation system on video views. *Proceedings of the 10th ACM SIGCOMM conference on Internet measurement* (2010), 404–410.