

NBER WORKING PAPER SERIES

USING INFORMATION TO CURB RACIAL DISCRIMINATION

Christopher R. Knittel
Donald MacKenzie
Michiko Namazu
Bora Ozaltun
Dan Svirsky
Stephen Zoepf

Working Paper 33118
<http://www.nber.org/papers/w33118>

NATIONAL BUREAU OF ECONOMIC RESEARCH
1050 Massachusetts Avenue
Cambridge, MA 02138
November 2024

As is customary in economics, names are listed alphabetically. Uber’s research team hosted this study, implemented the intervention, and gathered data. Under the agreement granted the authors, Uber has the right to review the paper “to identify any disclosure of Confidential Information to request its removal” and to “determine whether a patent application or other intellectual property protection should be sought prior to publication,” but not to dispute or influence the findings or conclusions of the paper. The views expressed herein are those of the authors and do not necessarily reflect the views of the National Bureau of Economic Research or Uber Technologies, Inc. Michiko Namazu and Dan Svirsky were employees and shareholders of Uber Technologies before, during, and after the writing of this paper.

NBER working papers are circulated for discussion and comment purposes. They have not been peer-reviewed or been subject to the review by the NBER Board of Directors that accompanies official NBER publications.

© 2024 by Christopher R. Knittel, Donald MacKenzie, Michiko Namazu, Bora Ozaltun, Dan Svirsky, and Stephen Zoepf. All rights reserved. Short sections of text, not to exceed two paragraphs, may be quoted without explicit permission provided that full credit, including © notice, is given to the source.

Using Information to Curb Racial Discrimination

Christopher R. Knittel, Donald MacKenzie, Michiko Namazu, Bora Ozaltun, Dan Svirsky,
and Stephen Zoepf

NBER Working Paper No. 33118

November 2024

JEL No. D83, J71, L92

ABSTRACT

We test whether more information about customers decreases racial bias. Our setting is the market for shared mobility services. Prior work by Ge et al. (2020) found that Uber drivers are two times more likely to cancel a ride if the passenger's name is one used predominantly by African Americans. In a randomized control trial, we test whether two alterations to the Uber platform app reduce racial discrimination. Within the standard Uber app, drivers see only the passenger's rating before accepting a ride. Once they accept the ride, they see the name of the passenger. In the first intervention, we increased the size of the font of the rating to draw attention to the quality of the passenger. In the second intervention, the passenger's name appears from the beginning. Using the control group observations, we confirmed that the more likely African Americans were to use a name, the more likely a driver cancels the ride. However, increasing the font size of the passenger's rating eliminates this racial bias. In contrast, we do not find much evidence that showing the name on the initial screen reduces or increases cancellation rates.

Christopher R. Knittel
MIT Sloan School of Management
100 Main Street, E62-513
Cambridge, MA 02142
and NBER
knittel@mit.edu

Bora Ozaltun
Department of Agricultural & Resource Economics
University of California, Berkeley
207 Giannini Hall
Berkeley, CA 94720
ozaltun@berkeley.edu

Donald MacKenzie
Proteus LLC
donaldwarrenmackenzie@gmail.com

Dan Svirsky
Uber Technologies, Inc.
dsvirsky@uber.com

Michiko Namazu
Uber Technologies, Inc.
nmichiko@uber.com

Stephen Zoepf
U.S. Department of Transportation
szoepf@mac.com

A randomized controlled trials registry entry is available at
<https://www.socialscienceregistry.org/trials/14658>

1 Introduction

Over the past decade, Transportation Network Company (TNC) services such as UberX have grown rapidly and, in many areas, are being used, formally or informally, to augment, expand or replace existing public and private transportation services. For example, transit agencies in Dallas, TX, Memphis, TN, and Durham, NC, have integrated fare payments with TNCs. Agencies in locations from Seattle to Orlando to Cape Cod have tested partnerships in which TNCs provide first- or last-mile connections to fixed-route transit services. As the importance of TNC services grows, so do concerns that they are delivering equal access to diverse communities and individuals.

Inequalities in access to TNC services may emerge from one or more behaviors by TNC drivers, including: (1) drivers choosing to avoid specific neighborhoods or to switch off the TNC app in those neighborhoods, leading to longer waiting times; (2) drivers declining to accept trip requests from certain types of passengers; (3) drivers accepting, but then canceling, a pickup for certain types of passengers; (4) drivers choosing non-ideal routes based on the same factors, increasing costs and/or travel time; or (5) drivers leaving low ratings for passengers based on factors such as race, gender, or perceived socioeconomic status, potentially hurting the traveler’s ability to book rides in the future.

Early research suggests a mixed record of such discrimination among drivers on major TNC platforms. On the one hand, research in Seattle, WA found no evidence of longer waiting times in neighborhoods with more minorities or lower average incomes (Hughes and MacKenzie, 2016), and drivers who completed trips gave virtually identical star ratings to black and white travelers (Ge et al., 2020). Academic research and popular press reports also suggest that Uber has improved access in underserved neighborhoods in New York (Hu, 2017) and Los Angeles Brown (2018). On the other hand, (Ge et al., 2020) used randomized field experiments to conclude that, within the Uber platform, a driver is twice as likely to cancel a ride if the passenger’s name is one used predominantly by African Americans.

Mitigating these inequalities in TNC service delivery is complicated because TNC drivers on the most popular platforms (e.g., Uber and Lyft) are independent contractors, so the platform operators have limited ability to dictate their behavior. Nevertheless, platform operators may be able to reduce discrimination—or the impacts of discrimination—by adjusting the timing, content, and presentation of passenger information delivered to drivers. On the Uber platform, racial discrimination manifests through cancellations because drivers do not see the passenger’s name until after they accept a trip request; the initial screen only shows the passenger’s rating. In contrast, Lyft drivers see the passenger’s name on the trip request screen.

Our focus in this work is on testing whether modifying the delivery of passenger information to drivers on the Uber platform reduces racial disparities in cancellation rates. We implement a randomized control trial to evaluate two interventions. The first increases the prominence of the passenger’s rating before a driver accepts the ride by increasing its font size. We refer to this as the “Rating-Font Treatment.” The second replicates the Lyft platform by showing the passenger’s name before the driver accepts the ride. We refer to this as the “Name Treatment.” Randomization

is across drivers; 1,200 drivers were assigned to each of three conditions in Boston, Los Angeles, and Seattle: a control arm or one of the two treatment arms for one month.

We first replicate the results of (Ge et al., 2020) by estimating how cancellations are affected by names within the control group. Using data from voting rolls in Georgia and North Carolina, we measure the share of people with a given first name who are African American. We then estimate how the probability of cancellation moves with the share of African Americans among people with the passenger’s name. Consistent with (Ge et al., 2020), we find that going from a name that is not predominantly used by African Americans to one that is exclusively used by African Americans leads to a 0.93 percentage point increase in the likelihood of cancellation – 15.6 percent of the average cancellation rate of 5.97 percent. We also estimate a piece-wise linear model and find that the cancellations increase sharply once the share of African Americans among people using that name reaches 50 percent.

Next, we estimate how the passenger’s name affects cancellations under our two interventions. We find that increasing the font size of the passengers’ ratings reduces the effect of the name on cancellations by over 71.7 percent. This reduction is statistically significant at the 5-percent level, and the cumulative impact of the passengers’ names on cancellations is no longer statistically significant. We do not find strong evidence that showing the passengers’ names before drivers accept the ride reduces cancellation rates among African Americans. The point estimate of the interaction between the treatment dummy suggests a 30.4 percent reduction in the effect of the passengers’ name on cancellations, but the point estimate is noisily estimated. These results suggest that platforms can mitigate discrimination by highlighting relevant information about those groups that are the target of discrimination.

We also investigate whether the interventions led to changes in where the drivers focused their efforts. To do this, we estimate how the interventions changed the average share of African-American residents in the census tracts where drivers pick up passengers. If the interventions led to drivers trying to avoid African American passengers, we would expect that the average share of African Americans in the locations where the driver picked up passengers would fall due to the intervention. We find, in contrast, that the average share of African Americans living in the pickup locations *increased* as a result of the rating-font intervention.

Previous literature

There is a large body of research on statistical discrimination spanning many fields. More specifically, our work looks at discrimination through different compositions of names by race. Fryer and Levitt (2004) discusses the differences in names used in black and white communities after the 1970s. Bertrand and Mullainathan (2004) studies the result of distinctive names and finds significant racial discrimination in callbacks to job applications.

Our intervention also contributes to the literature examining how information changes can impact discrimination. Goldin and Rouse (2000) uses the adoption of “blind” auditions by symphony orchestras for hiring as a mechanism to study discrimination that resulted in seeing/not seeing

a candidate. There has recently been a large literature on “Ban the Box” (BTB) policies and their impacts. Most significantly, Agan and Starr (2018) looks at the impact of BTB policies in New Jersey and New York City, finding that the policy has increased racial discrimination in job application callbacks.

There is also evidence of discrimination in the peer economy. Pope and Sydnor (2011) shows loan requests are less successful when associated with African-American applicants. Edelman and Luca (2014) finds discrimination towards Black landlords on Airbnb. Edelman et al. (2017) studies guests on Airbnb, finding guests with distinctively African American names are 16% less likely to be accepted.

2 Experimental Design

Our experiment ran in Boston, Los Angeles, and Seattle from September 2018 to November 2018. Before describing the treatments, it is important to describe the standard screen Uber drivers see when offered a fare. The left panel of Figure 1 shows a screenshot of what an Uber driver would see in 2018 when a trip is requested. The Figure shows that the driver sees various information about the potential trip. First, they see a map of the passenger’s location and the fastest route to the passenger. The text in the middle reports the type of trip it is (UberX in this case), the rating of the passenger (4.92), the time and distance to the passenger (3 minutes and 1.2 miles, respectively), and the surge price inflater of the trip (1.6x). Notably, the passenger’s name is not listed. For this reason, a cancellation is required by the driver to act on their bias.

We worked with Uber to implement two treatments. The first treatment increased the font size and color of the passenger’s rating. This is represented in the middle panel of Figure 1. Everything else about the initial screen remained unchanged. The second treatment added the passenger’s name to the initial screen (right panel), again keeping everything else the same. Randomization took place at the driver level. In each city, 3,600 drivers were included in the study, with 1,200 in the control group and 1,200 in each treatment group. Unlike in Ge et al. (2020), we did not randomize passenger names. We discuss the implications of this below.

Our main question is how the passenger’s name affects the likelihood that a driver will accept the ride and then cancel before pickup. Ge et al. (2020) used a set of 16 names, matching one name that Caucasians predominantly use with a name predominantly used by African Americans. In this experiment, we assign names on a more continuous basis. In particular, we use voting roll data from Georgia and North Carolina that report the voter’s name and race to calculate the percentage of people with that name: African American, Hispanic, and other.¹

We collected data on each of the requests offered to the driver. We observe the time, date, and Census tract of the request. We observe how long, in seconds, it took the driver to respond to

¹Using the combination of these two data sets, we are able to match $\approx 95\%$ of the passenger names. We estimate a series of robustness checks to see if different frequency thresholds for certain names impact the main results and find that the conclusions don’t change and the point estimates are statistically indistinguishable. See Table A6. Unlike Georgia, North Carolina’s data contains data on both race and ethnicity. For individuals that identify as Hispanic as an ethnicity and a certain race, we combine these by weighting both by 0.5 instead of 1.

the request, whether the driver accepted or declined, whether the passenger canceled the request, whether the driver arrived at the pickup location, and if so, how long it took them. We also know whether surge pricing was in place, the distance to the pickup location, the passenger’s rating, and the passenger’s reported first name. For all drivers, we observe requests for one month before treatments are implemented. Across treatment and control and the three cities, we have 1,381,215 trip requests. Of these trip requests, 1,114,509 of them were accepted by the driver.

3 Empirical Methods

Our basic empirical approach is to understand how passenger names impact the probability that a driver accepts the ride and then cancels before pickup and how this relationship is affected by the two treatments. This motivates the following difference-in-difference specification of the impact of the interventions on whether a cancellation takes place by driver d , for passenger with name i , in city c , and time t :

$$\begin{aligned} Cancel_{dict} = & \beta_1 AA_i + \theta_1 T_{dt}^{rating} AA_i + \theta_2 T_{dt}^{name} AA_i + \\ & \alpha_1 T_{dt}^{rating} + \alpha_2 T_{dt}^{name} + \sum_z \delta_z X_z + \eta_d + \eta_t + e_{dict}, \end{aligned} \quad (1)$$

where AA_i is the percentage of people with the given name i that are African American, T^x is one if the driver is in treatment arm x and time t is during the treatment period, X_z are a series of control variables, η_d are driver fixed effects, and η_t are time fixed effects. While X_z is not needed to obtain unbiased estimates of the treatment effects, we include them to increase statistical precision.² We also present results allowing for a piece-wise linear relationship between the name share and cancellations. We define five knots, the first three based on quartiles and the last two to get finer readings on the later half of the distribution.

To test whether our data are consistent with the cancellation effects of African American-sounding names found in Ge et al. (2020), we estimate how the probability of cancellation moves with X_i in the control group observations (and those rides that occurred before the treatments being implemented).

It is important to note that while drivers are randomized into control or one of the two treatments, we do not randomly assign passengers. This can bias the relationship between a passenger’s name and the probability of cancellation. For example, suppose drivers self-select such that those with stronger racial bias focus their driving on areas with a larger Caucasian population. In that case, the relationship between passenger names and cancellations will be muted, at least in terms of the bias reflected in whether race affects cancellations. While one could argue that the relationship of interest is the reduced form relationship after allowing for self-selection, fixed driver effects will account for this type of selection.

²We include passenger rating deciles, hour of the day indicators, day of the week indicators, an indicator variable for surge time periods, and the share of people with the passenger’s name that are Hispanic (HP). We cluster standard errors at the city-day level.

A final set of analyses looks to see if the treatments affect where drivers focus their efforts. To do this, we estimate how the treatments affected the characteristics of the Census tracts where drivers pick up passengers. In particular, we regress the share of African Americans living in the Census tract for a given completed ride. If drivers shifted their focus to areas with fewer African American residents in response to the treatment, then we would find a negative effect of the treatment on this outcome variable. Because we are interested in the driver’s overall behavior, we do not interact *AA* with the treatments in these specifications; we are only interested in the coefficient associated with the treatment indicator variables.

4 Results

We first test that our samples across treatment and control are balanced. Table 1 reports the means of the variable capturing the share of people with a given name that is African American, the average share of African Americans living in the Census tract of the pickup location, the average cancellation rate, the proportion of rides that are rideshare, and the share of passengers with a rider rating above the median. The last two columns report the p-value from a regression of the variable on a dummy variable for each treatment arm. The table shows balance in all of the variables except for the share of riders with a rating above the median. Here the p-values are 0.16 and 0.17 for the Rating-Font and Name Treatment arms, respectively.

Table 2 reports our baseline results. The first column presents the results from the difference-in-difference specification described in (1). The first thing to note is that we replicate the results in Ge et al. (2020) that African Americans are more likely to face cancellation. A passenger with a name solely used by African Americans faces a 0.93 percentage point increase in the likelihood of a cancellation. This is 15.6 percent of the average cancellation rate of 5.97 percent. We do not find evidence of discrimination against Hispanics.

The second column of Table 2 adds our treatment interactions. The rating-font treatment reduces the impact of *AA* on cancellations by 0.0066. This represents 71.7 percent of the *AA* coefficient. In addition, the sum of the net effect of the African-American share variable is no longer statistically significant. The second treatment—showing the passenger’s name before ride acceptance—also reduces the impact of *AA* on cancellations, but the point estimate is not statistically significant. The point estimate suggests a 30.4 percent reduction. This is somewhat surprising since there is no new information available to the driver upon acceptance; therefore, one might expect that showing the name ahead of time would completely eliminate the effect of *AA* on cancellations. One potential explanation is that drivers are accustomed to not seeing the name prior to accepting a ride and, therefore, do not look for the name until after acceptance. We also note that the confidence interval for *AA · NameTreatment* encompasses a range that would nearly eliminate the effect of *AA*.

In the third column, we investigate whether drivers altered where they chose to pick up passengers in response to the treatments. The concern here is that while treatment may eliminate

cancellations based on passenger names, if it leads to drivers avoiding certain areas then the welfare implications of such a treatment are unknown. As noted, to measure this, we look at the average share of African Americans living in each census tract for which a driver completed a drive. We find that drivers in the rating-font treatment group target areas with a *higher* share of African Americans compared to the control group. Therefore, drivers did not seem to avoid areas with a high share of African Americans. Indeed, they increased their pickups in these areas.

Figure 2 panel B plots the estimates from the piece-wise linear specifications for the control group observations. Interestingly, we find that the impact of the name is zero until the share of people with the name exceeds 50 percent. The coefficient associated with *AA* becomes even larger when the share of people with the name that is African American exceeds 70 percent. The coefficient dips for names for which over 90 percent of those with that name are African American. One possible explanation for the dip is that distinctly African American names in our data are also rarer, reducing the opportunity for discrimination. The overall results suggest that drivers do not discriminate based on the name until the share of people with that name who are African American exceeds 50 percent and that the greater the share, the more drivers discriminate based on the name.

We next allow the treatment to alter the shape of this piece-wise linear function. Figure 2 panels C and D plot the shape of this curve for the control drivers as well as the drivers in the two treatment groups. For the latter, this represents the sum of the *AA* and *AA · Treatment* variables. Panel C shows the results for the Rating-Font Treatment, while panel D shows the results of the Name Treatment. Confidence intervals are also included. The net effect for both treatments remains relatively flat when the share of people with the name that is African American is less than 50 percent. For the Rating-Font Treatment, the net effect of a change in the percentage of riders with a name that is African American (the slope of the piece-wise linear function) is similar over the range of 0.45 to 0.90. As with the control observations, the results are counter-intuitive and noisy when the share is above 90 percent.

The impacts of the Name Treatment are mixed. The slope of the piece-wise linear function increases when the share of riders with that name, African American (the slope of the piece-wise linear function), is between 0.45 and 0.70, but the function ends at a lower level given the lower starting point. Furthermore, the slope is flatter, compared to the control group, over the range of 0.70 to 0.90. This provides some evidence that the Name Treatment reduces cancellation when the share of riders with a given name that is African American is high.

Table 3 also reports the results for our basic treatment effect specification for each of the cities. All three cities exhibit a higher rate of cancellations the larger the share of African Americans using the given name. The online Appendix Table A1 includes the specifications using only the control group observations. The coefficients are similar across cities. The coefficient for Boston is 0.0086, and they are 0.0087 and 0.0119 for Los Angeles and Seattle, respectively. Interestingly, we see some evidence of discrimination against Hispanics in Boston. We note, however, that when we estimate the impact using only control group observations, the coefficient associated with Hispanics gets

smaller and noisier.

The impact of the Rating-Font treatment is much smaller in Boston compared to the other cities. In Los Angeles and Seattle, the Rating-Font treatment effect is nearly the size of the coefficient associated with *AA*. In Los Angeles, the treatment estimate is larger than the *AA* coefficient, while in Seattle, it is 82.1 percent of the *AA* coefficient. In Boston, however, it is very close to zero, although noisily estimated. In all cities, the coefficients associated with the Name treatment are smaller than the other treatment effects and not precisely estimated. This suggests that there may be significant treatment effect heterogeneity across locations. Finally, we find that the Rating-font treatment led drivers to increase the number of trips in African-American census tracts in both Boston and Seattle.

4.1 Alternative Specifications

We assess robustness across several dimensions. Ge et al. (2020) found that the impact of the rider’s name was larger in neighborhoods where the share of African Americans was small. This suggests that the neighborhood’s racial composition may be an important indicator of whether discrimination based on the name takes place. We include Census tract fixed effects in Appendix Table A2 in the specification. This allows us to estimate the impact of *AA* and the treatments within the census tract. We find that the impact of the name falls from 0.0092 in the specification without census tract effects to 0.0067 but remains statistically significant. We also find that the Rating-Font Treatment eliminates a greater share of discrimination (83.6 percent), although statistical significance wanes somewhat. The coefficient associated with the Name Treatment falls to near zero.

Our main specifications focus on cases where the driver explicitly cancels the ride. Drivers face negative repercussions when they have very high cancellation shares, so they may act strategically to get the rider to cancel the ride instead. This may take the form of simply ignoring the ride until the rider cancels. To account for this, we define the cancellation variable to include cases where the rider cancels after 15 minutes have elapsed since the driver accepted the ride.³ These results are reported in Appendix Table A3 and are consistent with our base results.

We also attempt to take out driver cancellations that would seem not to be motivated by discrimination. In particular, we omit cases where the driver cancels after reaching the location but cannot find the rider for an extended time. We define this time period as 3 minutes for Uber pool rides and 5 minutes for UberX rides⁴. These results are reported in Appendix Table A4. The general pattern of the coefficients is the same; however, the coefficients associated with the treatment effects are smaller and more noisily estimated.

We note again that drivers may respond to the Name Treatment by never accepting a ride in the first place. To study this, we define the dependent variable to be whether the driver accepts

³This adds 2,372 cancellations to the dependent variable, increasing the number of canceled rides from 66,500 to 68,872.

⁴We set different cutoffs since the enforced wait times by Uber pool versus UberX differ.

a ride and then cancels or whether the driver never accepts the ride in the first place. To test for this, Appendix Table A5 defines the dependent variable to one if the driver never accepts the request. Here, we find that the impact of the name increases from 0.0092 to 0.0130. Interpreting the Rating-font group estimates in this table isn't straightforward as we would expect no effect for drivers accepting rides (since they don't have the name information) while our main results show that there is a treatment effect for driver cancellation after acceptance. The very small and insignificant point estimate in Table A5 is in line with this interpretation. However, the Name treatment group results suggest that providing the name of the rider before acceptance does not have an impact on the discrimination from drivers. The point estimate is small and insignificant.

5 Conclusions

(Ge et al., 2020) found using a randomized field experiment in Boston that Uber drivers discriminate against riders with names that are associated with African Americans. By randomizing the names of passengers, they find that cancellations doubled when research assistants used a name for which a high share of people with that name is African American. Working directly with Uber, we implement a randomized control trial in an attempt to curb this discrimination. We test the effectiveness of two treatments. The first increases the prominence of the passenger's rating prior to a driver accepting the ride by increasing its font size. The second replicates the Lyft platform by showing the passenger's name prior to the driver accepting the ride.

Within our data, we first document the results in (Ge et al., 2020) by estimating how cancellations are affected by names within the control group. We find that going from a name that African Americans do not use to a name where all of the people with that name in our voting rolls are African American increases the likelihood of cancellation by 19.2 percent. We also estimate a piece-wise linear model and find that the cancellations increase precipitously once the share of people using that name that is African American reaches 50 percent.

Next, we estimate how passengers' names affect cancellations under our two interventions. We find that increasing the font size of the passengers' ratings reduces the effect of the name on cancellations by over 70 percent. The interaction of the treatment dummy variable with the name variable is statistically significant at the 5-percent level, and the cumulative effect of the passengers' names on cancellations is no longer statistically significant. We do not find strong evidence that showing the passengers' names prior to accepting the ride reduces cancellation rates among African Americans. The point estimate of the interaction between the treatment dummy suggests a 30-percent reduction in the effect of the passengers' names on cancellations, but the point estimate is noisily estimated.

Our findings carry significant business and policy implications for Transportation Network Companies (TNCs), and the sharing economy more generally. First, the effectiveness of providing tangible measures of rider quality in reducing the impact of passenger names on cancellations suggests a practical avenue for platforms to combat discrimination. Implementing similar modifications in the

presentation of passenger information, such as highlighting relevant details like ratings, could be a strategic move for platform companies aiming to enhance equal access and diminish disparities. However, the nuanced relationship observed with the “Name Treatment” emphasizes the need for careful consideration of interventions, urging platform companies to adopt evidence-based strategies. From a policy standpoint, regulators and policymakers can leverage these insights to enact industry-wide measures that prioritize anti-discrimination, while ensuring that changes are guided by evidence. As platform industries continue to evolve, our study underscores the importance of proactive engagement with these issues, urging stakeholders to collaborate in shaping policies that foster inclusivity and rectify disparities within the evolving landscape of transportation services.

References

- AGAN, A. AND S. STARR (2018): “Ban the box, criminal records, and racial discrimination: A field experiment,” *The Quarterly Journal of Economics*, 133, 191–235.
- BERTRAND, M. AND S. MULLAINATHAN (2004): “Are Emily and Greg More Employable Than Lakisha and Jamal? A Field Experiment on Labor Market Discrimination,” *The American Economic Review*, 94, 991–1013.
- BROWN, A. E. (2018): “Ridehail revolution: Ridehail travel and equity in Los Angeles,” Ph.D. thesis, UCLA.
- EDELMAN, B., M. LUCA, AND D. SVIRSKY (2017): “Racial discrimination in the sharing economy: Evidence from a field experiment,” *American economic journal: applied economics*, 9, 1–22.
- EDELMAN, B. G. AND M. LUCA (2014): “Digital discrimination: The case of Airbnb. com,” *Harvard Business School NOM Unit Working Paper*.
- FRYER, R. G. AND S. D. LEVITT (2004): “The Causes and Consequences of Distinctively Black Names,” *The Quarterly Journal of Economics*, 119, 767–805.
- GE, Y., C. R. KNITTEL, D. MACKENZIE, AND S. ZOEPEF (2020): “Racial discrimination in transportation network companies,” *Journal of Public Economics*, 190, 104205.
- GOLDIN, C. AND C. ROUSE (2000): “Orchestrating impartiality: The impact of “blind” auditions on female musicians,” *American economic review*, 90, 715–741.
- HU, W. (2017): “Uber, Surging Outside Manhattan, Tops Taxis in New York City,” *New York Times*, October 12th.
- HUGHES, R. AND D. MACKENZIE (2016): “Transportation network company wait times in Greater Seattle, and relationship to socioeconomic indicators,” *Journal of Transport Geography*, 56, 36–44.
- POPE, D. G. AND J. R. SYDNOR (2011): “Implementing anti-discrimination policies in statistical profiling models,” *American Economic Journal: Economic Policy*, 3, 206–31.

Tables

Table 1: Average Characteristics and Treatment-Control Balance

	(1)	(2)	(3)	(1)	(2)
	Treatment (rating)	Treatment (name)	Control	Treatment p-val (rating)	Treatment p-val (name)
AA	0.19	0.19	0.19	0.38	0.62
Census tract - AA	0.09	0.09	0.08	0.00	0.47
1(Driver cancels)	0.05	0.04	0.05	0.17	0.40
1(Ride share)	0.23	0.23	0.24	0.18	0.11
1(Above median rating)	0.53	0.53	0.48	0.16	0.17

Notes: This table provides average values from different variables of the dataset in the first three columns. The last two columns report p-values of regressing the row variable on the corresponding treatment indicator with fixed effects for hour, if the ride is a surge ride, day of week, and driver id. Finally, the p-values are obtained from standard errors that are clustered by city.

Table 2: Main results

	(1)	(2)	(3)
	1(Driver cancels)	1(Driver cancels)	Census tract: AA
AA	0.0093*** (0.0015)	0.0092*** (0.0015)	
HP	-0.0001 (0.0024)	0.0009 (0.0019)	
1(Rating-font treatment):AA		-0.0066** (0.0031)	
1(Name treatment):AA		-0.0028 (0.0031)	
1(Name treatment)		-0.0006 (0.0015)	0.0007 (0.0005)
1(Rating-font treatment)		-0.0007 (0.0014)	0.0015** (0.0006)
Observations	706,493	1,114,509	1,114,509

*Notes: This table provides the main results for the study. The first column contains the estimates run on the control group, whereas the second and third column contain the estimates run on the whole dataset. Reminder that * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.*

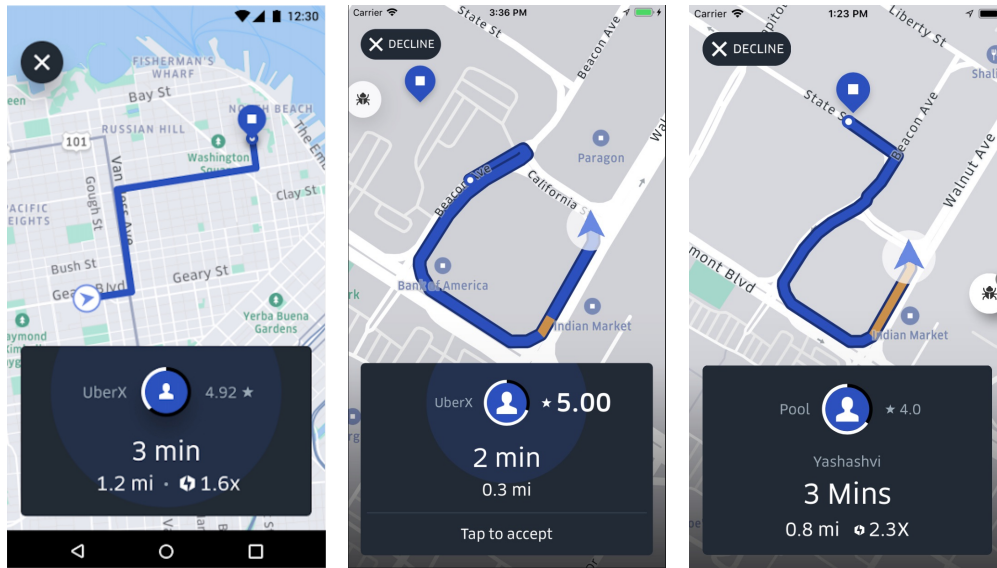
Table 3: Results in each city

	(1)	(2)	(3)	(4)	(5)	(6)
	Boston	Boston	Los Angeles	Los Angeles	Seattle	Seattle
	1(Driver cancels)	CT: AA	1(Driver cancels)	CT: AA	1(Driver cancels)	CT: AA
AA	0.0090** (0.0035)		0.0086*** (0.0016)		0.0117*** (0.0028)	
HP	0.0084* (0.0038)		-0.0023 (0.0016)		0.0001 (0.0032)	
$T^{rating}.AA$	-0.0007 (0.0068)		-0.0102** (0.0036)		-0.0096* (0.0047)	
$T^{name}.AA$	-0.0051 (0.0073)		0.0009 (0.0036)		-0.0054 (0.0044)	
T^{name}	0.0050** (0.0017)	0.0010 (0.0009)	-0.0020 (0.0013)	0.0006 (0.0007)	-0.0063* (0.0027)	0.0003 (0.0009)
T^{rating}	0.0005 (0.0028)	0.0028** (0.0011)	0.0002 (0.0016)	0.0007 (0.0017)	-0.0030 (0.0023)	0.0011* (0.0006)
Observations	412,813	412,813	348,127	348,127	353,569	353,569

Notes: This table provides the city level results for the study. Columns (1) and (2) are the results from Boston, (3) and (4) are the results from Los Angeles, and (5) and (6) are the results from Seattle. Note that the two columns for each city have the same estimation setup as the second and third column of Table 2, but using the city-level subsets. Additionally, some of the variable labels have been shortened for brevity: (1) T^{name} and T^{rating} correspond to the 1(Name treatment) and 1(Rating-font treatment) and (2) CT: AA corresponds to census tract level data on ratio of African American residents. Reminder that * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.

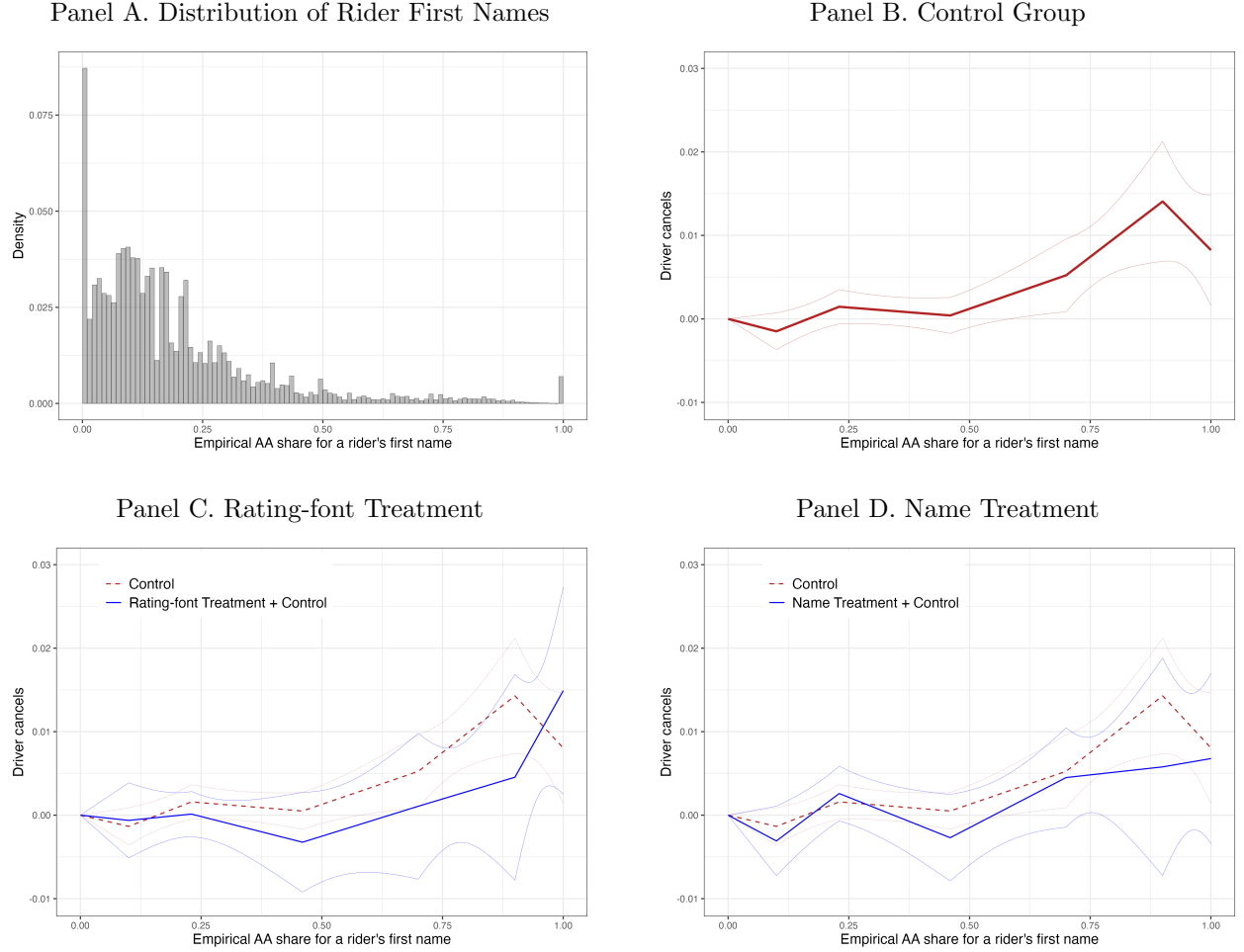
Figures

Figure 1: Control, Rating-Font Treatment, and Name Treatment Screenshots



Notes: This figure provides screenshots for the three different groups within the experiment. The left most screenshot is the control group. At the time of a ride request the location of the driver's current location and the location of the passenger are shown in the map. The legend includes the passenger's rating, the estimated time to reach the passenger, distance to the passenger, and any surge multiplier that exists. Notably, the name of the passenger does not appear. Our first treatment arm increases the size and brightness of the passenger's rating, keeping all other information the same. The second treatment provides the name of the passenger, keeping everything else the same.

Figure 2: Heterogeneous Treatment Effects



Notes: The top left panel plots AA_i for each individual in our dataset (where AA_i is the percentage of people with the given name i that are African American). The top right panel is the of piece-wise linear relationship between cancellations and the share of people that is African American for a given name among the control-group observations. The next two panels plot the second panel (dotted) along with this same relationship among the two treatment groups.

A Internet Appendix

A.1 Tables

Table A1: Results in each city: Control group

	(1) Boston 1(Driver cancels)	(2) Los Angeles 1(Driver cancels)	(3) Seattle 1(Driver cancels)
AA	0.0086** (0.0035)	0.0087*** (0.0014)	0.0119*** (0.0027)
HP	0.0057 (0.0054)	−0.0024 (0.0025)	−0.0001 (0.0050)
Observations	261,429	211,089	233,975

Notes: This table provides the city level results for the control group. Each column contains the estimates run on the control group. Reminder that * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.

Table A2: Results with Alternative Specifications: Adding Census Tract Fixed Effects

	(1) 1(Driver cancels)	(2) 1(Driver cancels)	(3) Census tract: AA
1(Name treatment)	0.0063*** (0.0016)	0.0067*** (0.0016)	
1(Rating-font treatment)	−0.00003 (0.0025)	0.0010 (0.0018)	
AA		−0.0056* (0.0031)	
HP		−0.0017 (0.0033)	
1(Rating-font treatment):AA		−0.0009 (0.0014)	0.0000 (0.0000)
1(Name treatment):AA		−0.0003 (0.0013)	−0.0000 (0.0000)
Observations	706,493	1,114,509	1,114,509

Notes: This table provides the results for an alternative specification of the main study. The difference between the main specification and the results in this table is that this table adds Census Tract fixed effects. The first column contains the estimates run on the control group, whereas the second and third column contain the estimates run on the whole dataset. Reminder that * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.

Table A3: Results with Alternative Specifications: Adding rider cancellations after 15 minutes

	(1) 1(Driver cancels) or 1(Rider cancels after 15 minutes)	(2) 1(Driver cancels) or 1(Rider cancels after 15 minutes)
1(Name treatment)	0.0094*** (0.0016)	0.0095*** (0.0016)
1(Rating-font treatment)	0.0002 (0.0023)	0.0012 (0.0019)
AA		−0.0066* (0.0034)
HP		−0.0036 (0.0030)
1(Rating-font treatment):AA		−0.0009 (0.0015)
1(Name treatment):AA		−0.0009 (0.0015)
Observations	706,493	1,114,509

*Notes: This table provides the results for an alternative specification of the main study. The difference between the main specification and the results in this table is the dependent variable used. In this table, we include rider cancellations after 15 minutes have elapsed in order to account for driver's potentially acting strategically. The first column contains the estimates run on the control group whereas the second column contains the estimates run on the whole dataset. Reminder that * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.*

Table A4: Results with Alternative Specifications: Rider no show

	(1)	(2)
	1(Driver cancels) and 1(Rider no show)=0	1(Driver cancels) and 1(Rider no show)=0
1(Name treatment)	0.0043*** (0.0014)	0.0044*** (0.0014)
1(Rating-font treatment)	-0.0029 (0.0018)	-0.0020 (0.0016)
AA		-0.0030 (0.0030)
HP		-0.0024 (0.0024)
1(Rating-font treatment):AA		-0.0005 (0.0015)
1(Name treatment):AA		-0.0013 (0.0014)
Observations	706,493	1,114,509

*Notes: This table provides the results for an alternative specification of the main study. The difference between the main specification and the results in this table is the dependent variable used. In this table, we exclude driver cancellations where we suspect that the rider isn't showing up. We define this by setting limits on how long the driver waits for a rider. 3 minutes for ride-sharing and 5 minutes for UberX rides. The first column contains the estimates run on the control group whereas the second column contains the estimates run on the whole dataset. Reminder that * $p < 0.1$; ** $p < 0.05$; *** $p < 0.01$.*

Table A5: Results with Alternative Specifications: Driver doesn't accept

	(1) 1(Driver cancels) or 1(Driver doesn't accept)	(2) 1(Driver cancels) or 1(Driver doesn't accept)
1(Name treatment)	0.0134*** (0.0028)	0.0130*** (0.0028)
1(Rating-font treatment)	0.0046 (0.0045)	0.0062 (0.0038)
AA		−0.0006 (0.0042)
HP		−0.0002 (0.0055)
1(Rating-font treatment):AA		0.0047 (0.0031)
1(Name treatment):AA		−0.0009 (0.0039)
Observations	873,258	1,381,215

*Notes: This table provides the results for an alternative specification of the main study. The difference between the main specification and the results in this table is the dependent variable used. In this table, we include cases where the driver doesn't accept the offer in the first place. The first column contains the estimates run on the control group whereas the second column contains the estimates run on the whole dataset. Reminder that $*p < 0.1$; $**p < 0.05$; $***p < 0.01$.*

Table A6: Restricting the dataset to having at least 5 observations of the same name in the voter registration database

	(1) 1(Driver cancels)	(2) 1(Driver cancels)	(3) Census tract: AA
AA	0.0093*** (0.0019)	0.0093*** (0.0019)	
HP	−0.0002 (0.0024)	0.0003 (0.0019)	
1(Rating-font treatment):AA		−0.0090** (0.0032)	
1(Name treatment):AA		−0.0023 (0.0037)	
1(Name treatment)		−0.0008 (0.0014)	0.0007 (0.0006)
1(Rating-font treatment)		−0.0004 (0.0013)	0.0015** (0.0007)
Observations	677,202	1,067,927	1,067,927

*Notes: This table provides the results when the matching of names between voter registration and passenger names is restricted to names with at least 5 observations. The first column contains the estimates run on the control group, whereas the second and third column contain the estimates run on the whole dataset. Reminder that $*p < 0.1$; $**p < 0.05$; $***p < 0.01$.*