
Exploring the Impact of (Not) Changing Default Settings in Algorithmic Crime Mapping - A Case Study of Milwaukee, Wisconsin

MD Romael Haque

mdromael.haque@marquette.edu
Marquette University
Milwaukee, Wisconsin, USA

Katherine Weathington

katy.weathington@marquette.edu
Marquette University
Milwaukee, Wisconsin, USA

Shion Guha

shion.guha@marquette.edu
Marquette University
Milwaukee, Wisconsin, USA

ABSTRACT

Policing decisions, allocations and outcomes are determined by mapping historical crime data geospatially using popular algorithms. In this extended abstract, we present early results from a mixed-methods study of the practices, policies, and perceptions of algorithmic crime mapping in the city of Milwaukee, Wisconsin. We investigate this differential by visualizing potential demographic biases from publicly available crime data over 12 years (2005-2016) and conducting semi-structured interviews of 19 city stakeholders and provide future research directions from this study.

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

CSCW '19 Companion, November 9–13, 2019, Austin, TX, USA

© 2019 Copyright held by the owner/author(s).

ACM ISBN 978-1-4503-6692-2/19/11.

<https://doi.org/10.1145/3311957.3359500>

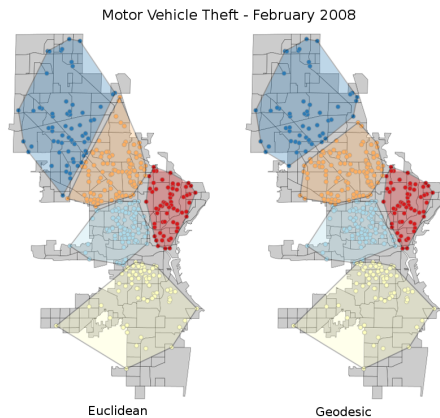


Figure 1: Comparison of Euclidean and Geodesic k-means clustering for Motor Vehicle Theft for February 2008

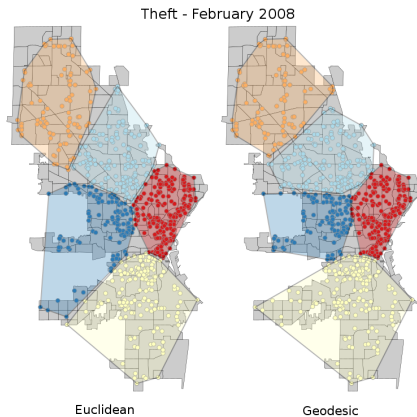


Figure 2: Comparison of Euclidean and Geodesic k-means clustering for Theft for February 2008

INTRODUCTION

Algorithms have become pervasive [11] in most facets of daily living. Recognizing the growing importance of algorithmic transparency debate, HCI/CSCW researchers have slowly started crafting a broad research agenda in this area including thinking about how data analysts engage in the act of analyzing data [13] and how experts, non-experts and subjects perceive data [1] to support such goals. One of the most common applications of algorithms [4, 14] is in the area of crime analysis. Crime analysis focuses on crime mapping, prediction and forecasting. Results are usually used to develop administrative policies that allocate policing resources to particular geographical areas or to focus on specific crimes. What effect could the combination of algorithmic opacity and knowledge have on the ethical mapping of crime as crime analysts grapple and interact with ever increasing and complex forms of data? Our research project is attempting to understand such practices and their potentially unanticipated future consequences through a human-centered lens.

In this extended abstract, we present some initial findings of our mixed methods study of the perceptions, practices and policies of algorithmic crime mapping in the city of Milwaukee, Wisconsin. We investigated publicly available crime data over a period of 12 years (2005-2016) and conduct a semi-structured interview study of 19 professional crime analysts and city stakeholders. Combining our methodological approaches, our initial exploration of the study suggests some theoretical implications such as default behaviors analysis among crime analysts.

DEFAULT SETTINGS IN ALGORITHMS

A default refers to predetermined parameters or settings that are being fixed by a computer program when a parameters or setting is not specified by the program user [15]. Past work has found that default policies can have a profound impact on users' final policies and their overall use of a system. For example, users tend not to change default calendar sharing settings [12], online social network privacy settings [2, 6, 9, 16], and even organ donation choices [8].

Clearly, how policymakers select the default has important implications. Policymakers often have to decide which of the available options to impose on individuals who fail to make a decision [3] as people perceive the default as indicating the recommended course of action. It is very important for the policymakers to be aware of the implied messages conveyed by their choice of default as the user might rationally decide to stick with this default if he or she adequately trusts the system [10].

METHODS

We started by interviewing two professional crime analysts to get an initial insights into algorithmic crime mapping practices. We used publicly available crime data about the city of Milwaukee for 12 years (2005-2016) as an empirical lens of investigation. We focused on the 'k-means' algorithm

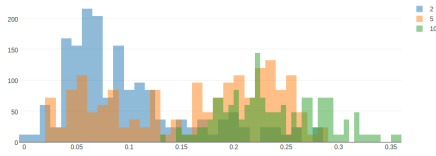


Figure 3: A histogram showing potential bias index (PBI) frequency

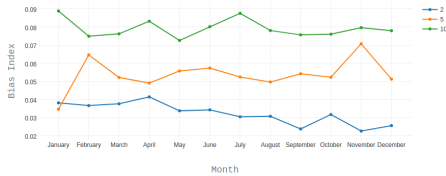


Figure 4: potential bias index (PBI) averages for each month for 2, 5 and 10 clusters

ALGORITHM 1: Potential Bias Index

Input: G : geodesic cluster

Input: E : list of unique euclidean clusters in G

Output: I : Potential Bias Index

$numGeodesicPoints \leftarrow getPointCount(G)$

$minorityRatio \leftarrow getMinorityRatio(G)$

$clusterScore \leftarrow 0$

// for each euclidean cluster found in G

foreach $e_i \in E$ **do**

 // for each point in geodesic cluster

$euclideanPoints \leftarrow 0$

$matches \leftarrow 0$

foreach $p_j \in e_i$ **do**

 // if euclidean point is in geodesic cluster

if $p_j \in G$ **then**

$matches \leftarrow matches + 1$

end

$numEuclideanPoints \leftarrow numEuclideanPoints + 1$

end

$score \leftarrow matches / numEuclideanPoints$

$weight \leftarrow matches / numGeodesicPoints$

$index \leftarrow score * weight$

$clusterScore \leftarrow clusterScore + index$

end

$dissimilarity \leftarrow 1 - clusterScore$

$potentialBiasIndex \leftarrow dissimilarity * minorityRatio$

return $potentialBiasIndex$

Figure 5: Potential Bias Index Algorithm

because its' flaws are intuitive to understand for the layperson. We restricted our analysis to four common crimes: robbery, simple assault, theft and and motor vehicle theft that are commonly mapped by analysts. We created visualizations of potential bias and used publicly available demographic information to create a Potential Bias Index (PBI) (Fig 5.) that we used as visual aids in the next round of interviews.

Then, we conducted follow-up interviews of 17 people. Eleven of them were professional crime analysts also working in the greater Milwaukee and Chicago metropolitan area. Six participants were local community organizers working to improve opportunities and reduce crime in the inner city. We adopted a grounded theory perspective [5] to our work. After multiple iteration of thematic analysis, initial high level themes have been emerged from the qualitative data.

INITIAL RESULTS & DISCUSSION

Deconstructing k-means for potential biases

Examining Lloyd's algorithm for k-means, we found two inflection points for potential human bias [7] i.e. (a) the initial selection of clusters and (b) the choice of the distance metric. Considering (a) (Fig. 3), in practice, values for both theft and motor vehicle theft ranged from 0 to a high of 0.36. The average potential bias for a given k ranged between 0.069 and 0.17 for theft and between 0.063 and 0.1706 for motor vehicle theft. In general, values of k greater than 4 produced an average bias value greater than or equal to .14, while values of k less than 4 produced values less than 0.1.

For theft, the gold standard of 5 clusters produced a low potential bias value of 0.0315 and a high value of 0.3099 with a mean of 0.1442 and standard deviation of 0.0562. Motor Vehicle Theft had a larger range with a low of 0.0180, a high 0.3495, a mean of 0.1457, and a standard deviation 0.0665. Theft exhibited lower standard deviation than motor vehicle theft, likely due to the higher number of data points (900 vs 400). But between both, when high potential bias values are produced, the associated clusterings typically featured two different configurations of the city center, while the clusters in the northern and southern ends of the city tended to be similar. This is likely due to the sparser nature of points on the city periphery, while the density of points toward the center of the city created more "unstable" initializations that result in high potential bias scores.

Considering (b) and looking at a given geodesic cluster, dissimilarity can increase in two ways. First, dissimilarity will increase when the number of unique euclidean clusters present increases. Geodesic cluster purity will decrease dissimilarity. Second, dissimilarity will increase if a small ratio of euclidean points are found inside the geodesic cluster compared to the number of points in the euclidean cluster. This dissimilarity score can be between 0 and 1. Zero means a geodesic cluster matches perfectly with a euclidean cluster. If a geodesic cluster contains small fractions of many different euclidean clusters, its score will approach 1. A visualization of this effect is presented in Figure 1 and 2.

"I didn't know what these distance things [metrics] are...I understand the Euclidean that...the calculation of the straight line because we learnt it in high school but I didn't know that there were other ways to calculate distance. I just point and click [on the GUI based crime analysis software that they use developed by a private third party]"- Jill (28, female, crime analyst)

"When I go to run the clusters [referring to k-means or other clustering methods], there are many other options on the menu but I don't know most of them so I just go with the default options on the menu... we were taught a basic idea of clustering but I didn't know that we could have so many different options. - John (37, male, crime analyst)"

"When I started the job, I was told that we always divide the city into five main divisions. There is the downtown cluster, the northshore cluster where all the rich folks live...you have the northwestern and southside clusters where there is a lot of gang activity and then the west side near the suburbs where a lot of people commute from." - Kevin (29, male, crime analyst)

"I am not sure how this [k-means algorithm] works. In school, we were always taught to think about applying the right tool for the right job but we weren't taught much about what's under the hood ...we were told that it [k-means] works very well for spatial data but we didn't learn much else." - Matthew (34, male, crime analyst)

Default behavior of Crime analysts

One of the main findings from our interviews is that, on the whole, crime analysts were unclear about the theoretical design and inner workings of the algorithms that they were using. Decisions made during data analysis were mostly supplemented with prior knowledge and existing mental models of the city.

All our analyst interviewees had masters degrees in criminology, crime analysis, sociology or public administration and had taken a few courses in applied statistics like Mathew. Some participants reported complete unfamiliarity with statistical distance metrics after we explained how k-means worked and displayed our visualizations like Jill. In this case, Jill does not change the default distance metric (Euclidean) that is provided in the software even though other options are present. Others point to a lack of transparency and clarity within the choices provided by the software that they use and a confusion in selecting appropriate options. This leads them to select default options. For instance, what John said in the given quote.

This refers to a general lack of transparency in how this third party software designs and implements the algorithms. When faced with a variegated menu of choices, the analysts select the one that is most familiar i.e. the default option. Taken together, this type of analysis is rule-based and path-bound[13]. It is natural to be paralyzed by a suite of potential options and then choose the most familiar one, however incorrect it might be under the given circumstances. However, when asked about how they decide to select the initial number of clusters, some participants responded that they depended on existing institutional knowledge about crime in Milwaukee. For instance, when asked about city-level clustering, Kevin referred to extant institutional knowledge that is in all likelihood, already biased.

Any subsequent analysis depends on this initial categorization that is dependent on institutional knowledge. Therefore, this type of analysis is based on situated decision making[13]. We observe here that while domain knowledge is very important, when combined together with what we learnt about the statistical (in)appropriateness of the actual process, there is a lot of potential for mis-classification and untoward policy making. Relatively few people request to switch from the default regardless of what the default is. Clearly, the default selected by policymakers has important implications.

CONCLUSION

We presented an exploratory analysis of the ways in which opacity and bias affects professional crime analysis by focusing on the practices, policies and perceptions around crime in Milwaukee, Wisconsin. We used publicly available data over a 12 year period (2005-2016) as well as interviews of 19 stakeholders (professional crime analysts and community organizers) to make our case. Moreover, our effort in involving multiple stakeholders to understand this issue is showed to be very illuminating especially in understanding practices of police departments around crime analysis.

REFERENCES

- [1] Eric P.S. Baumer, Xiaotong Xu, Christine Chu, Shion Guha, and Geri K. Gay. 2017. When Subjects Interpret the Data: Social Media Non-use As a Case for Adapting the Delphi Method to CSCW. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing (CSCW '17)*. ACM, New York, NY, USA, 1527–1543. <https://doi.org/10.1145/2998181.2998182>
- [2] Joseph Bonneau and Sören Preibusch. 2010. The Privacy Jungle: On the Market for Data Protection in Social Networks. In *Economics of Information Security and Privacy*, Tyler Moore, David Pym, and Christos Ioannidis (Eds.). Springer US, Boston, MA, 121–167.
- [3] Colin Camerer, Samuel Issacharoff, George Loewenstein, Ted O'Donoghue, and Matthew Rabin. 2003. *Regulation for Conservatives: Behavioral Economics and the Case for 'Asymmetric Paternalism'*. SSRN Scholarly Paper ID 399501. Social Science Research Network, Rochester, NY. <https://papers.ssrn.com/abstract=399501>
- [4] Hsinchun Chen, Wingyan Chung, Jennifer Jie Xu, Gang Wang, Yi Qin, and Michael Chau. 2004. Crime data mining: a general framework and some examples. *computer* 37, 4 (2004), 50–56.
- [5] Barney Glaser. 2017. *Discovery of Grounded Theory: Strategies for Qualitative Research*. Routledge. Google-Books-ID: GTMrDwAAQBAJ.
- [6] Ralph Gross and Alessandro Acquisti. 2005. Information Revelation and Privacy in Online Social Networks. In *Proceedings of the 2005 ACM Workshop on Privacy in the Electronic Society (WPES '05)*. ACM, New York, NY, USA, 71–80. <https://doi.org/10.1145/1102199.1102214>
- [7] J. A. Hartigan and M. A. Wong. 1979. Algorithm AS 136: A K-Means Clustering Algorithm. *Journal of the Royal Statistical Society. Series C (Applied Statistics)* 28, 1 (1979), 100–108. <https://doi.org/10.2307/2346830>
- [8] Eric J. Johnson and Daniel Goldstein. 2003. Do Defaults Save Lives? *Science* 302, 5649 (2003), 1338–1339. <https://doi.org/10.1126/science.1091721> arXiv:<https://science.sciencemag.org/content/302/5649/1338.full.pdf>
- [9] Kevin Lewis, Jason Kaufman, and Nicholas Christakis. 2008. The Taste for Privacy: An Analysis of College Student Privacy Settings in an Online Social Network. *Journal of Computer-Mediated Communication* 14, 1 (2008), 79–100. <https://doi.org/10.1111/j.1083-6101.2008.01432.x> arXiv:<https://onlinelibrary.wiley.com/doi/pdf/10.1111/j.1083-6101.2008.01432.x>
- [10] Craig R.M. McKenzie, Michael J. Liersch, and Stacey R. Finkelstein. 2006. Recommendations Implicit in Policy Defaults. *Psychological Science* 17, 5 (2006), 414–420. <https://doi.org/10.1111/j.1467-9280.2006.01721.x> arXiv:<https://doi.org/10.1111/j.1467-9280.2006.01721.x> PMID: 16683929.
- [11] Cathy O'Neil. 2016. *Weapons of math destruction: How big data increases inequality and threatens democracy*.
- [12] Leysia Palen. 1999. Social, Individual and Technological Issues for Groupware Calendar Systems. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems (CHI '99)*. ACM, New York, NY, USA, 17–24. <https://doi.org/10.1145/302979.302982>
- [13] Samir Passi, Steven Jackson, Phoebe Sengers, Almila Akdag Salah, Sally Wyatt, and Andrea Schornhorst. 2017. Data Vision: Learning to See Through Algorithmic Abstraction.. In *CSCW*. 2436–2447.
- [14] M.I. Pramanik, Raymond Y.K. Lau, Wei T. Yue, Yunming Ye, and Chunping Li. 2017. Big data analytics for security and criminal investigations: Big data analytics for security and criminal investigations. *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery* 7, 4 (July 2017), e1208. <https://doi.org/10.1002/widm.1208>
- [15] Margaret Rouse. 2005. default. Retrieved June 22, 2021 from <https://whatis.techtarget.com/definition/default>
- [16] Na Wang, Pamela Wisniewski, Heng Xu, and Jens Grossklags. 2014. Designing the Default Privacy Settings for Facebook Applications. In *Proceedings of the Companion Publication of the 17th ACM Conference on Computer Supported Cooperative Work & Social Computing (CSCW Companion '14)*. ACM, New York, NY, USA, 249–252. <https://doi.org/10.1145/2556420.2556495>