



Leveraging Simulation Data to Understand Bias in Predictive Models of Infectious Disease Spread

ANDREAS ZÜFLE, Computer Science, Emory University, Atlanta, United States

FLORA SALIM, School of Computer Science and Engineering, University of New South Wales, Sydney, Australia

TAYLOR ANDERSON, George Mason University, Fairfax, United States

MATTHEW SCOTCH, Arizona State University, Tempe, United States

LI XIONG, Emory University, Atlanta, United States

KACPER SOKOL, ETH Zurich, Zurich, Switzerland

HAO XUE, University of New South Wales, Sydney, Australia

RUOCHEN KONG, Computer Science, Emory University, Atlanta, United States

DAVID HESLOP, University of New South Wales, Sydney, Australia

HYE-YOUNG PAIK, University of New South Wales, Sydney, Australia

C. RAINA MACINTYRE, University of New South Wales, Sydney, Australia

The spread of infectious diseases is a highly complex spatiotemporal process, difficult to understand, predict, and effectively respond to. Machine learning and artificial intelligence (AI) have achieved impressive results in other learning and prediction tasks; however, while many AI solutions are developed for disease prediction, only a few of them are adopted by decision-makers to support policy interventions. Among several issues preventing their uptake, AI methods are known to amplify the bias in the data they are trained on. This is especially problematic for infectious disease models that typically leverage large, open, and inherently biased spatiotemporal data. These biases may propagate through the modeling pipeline to decision-making, resulting in inequitable policy interventions. Therefore, there is a need to gain an understanding of how the AI disease modeling pipeline can mitigate biased input data, in-processing models, and biased outputs. Specifically, our vision is to develop a large-scale micro-simulation of individuals from which human mobility, population,

This research is supported by the United States National Science Foundation under Grants No. 2302968, No. 2302969, and No. 2302970, titled “Collaborative Research: NSF-CSIRO: HCC: Small: Understanding Bias in AI Models for the Prediction of Infectious Disease Spread” as well as by the Australian Commonwealth Scientific and Industrial Research Organisation (CSIRO), and National Science Foundation under Grants No. 2125530 and No. 2041952.

Authors’ Contact Information: Andreas Züfle, Computer Science, Emory University, Atlanta, Georgia, United States; e-mail: azufle@emory.edu; Flora Salim, School of Computer Science and Engineering, University of New South Wales, Sydney, New South Wales, Australia; e-mail: flora.salim@unsw.edu.au; Taylor Anderson, George Mason University, Fairfax, Virginia, United States; e-mail: tander6@gmu.edu; Matthew Scotch, Arizona State University, Tempe, Arizona, United States; e-mail: Matthew.Scotch@asu.edu; Li Xiong, Emory University, Atlanta, Georgia, United States; e-mail: lxiong@emory.edu; Kacper Sokol, ETH Zurich, Zurich, Zürich, Switzerland; e-mail: ks1591@my.bristol.ac.uk; Hao Xue, University of New South Wales, Sydney, New South Wales, Australia; e-mail: hao.xue1@unsw.edu.au; Ruochen Kong, Computer Science, Emory University, Atlanta, Georgia, United States; e-mail: ruochen.kong@emory.edu; David Heslop, University of New South Wales, Sydney, New South Wales, Australia; e-mail: d.heslop@unsw.edu.au; Hye-Young Paik, University of New South Wales, Sydney, New South Wales, Australia; e-mail: h.paik@unsw.edu.au; C. Raina MacIntyre, University of New South Wales, Sydney, New South Wales, Australia; e-mail: r.macintyre@unsw.edu.au.

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

© 2024 Copyright held by the owner/author(s). Publication rights licensed to ACM.

ACM 2374-0353/2024/06-ART17

<https://doi.org/10.1145/3660631>

and disease ground-truth data can be obtained. From this complete dataset—which may not reflect the real world—we can sample and inject different types of bias. By using the sampled data in which bias is known (as it is given as the simulation parameter), we can explore how existing solutions for fairness in AI can mitigate and correct these biases and investigate novel AI fairness solutions. Achieving this vision would result in improved trust in such models for informing fair and equitable policy interventions.

CCS Concepts: • **Information systems** → **Spatial-temporal systems**; • **Applied computing** → *Health informatics*;

Additional Key Words and Phrases: Spatial epidemiology, fair and equitable AI, spatiotemporal prediction, epidemic forecasting

ACM Reference Format:

Andreas Züfle, Flora Salim, Taylor Anderson, Matthew Scotch, Li Xiong, Kacper Sokol, Hao Xue, Ruochen Kong, David Heslop, Hye-Young Paik, and C. Raina MacIntyre. 2024. Leveraging Simulation Data to Understand Bias in Predictive Models of Infectious Disease Spread. *ACM Trans. Spatial Algorithms Syst.* 10, 2, Article 17 (June 2024), 22 pages. <https://doi.org/10.1145/3660631>

1 INTRODUCTION

The 21st Century has been marked by major epidemics and pandemics caused by infectious diseases like sarbecoviruses (such as SARS-CoV-2), monkeypox, and influenza. Among the many strategies to manage these diseases, epidemiological models can advance our understanding of disease ecology and evolution. Popular modeling approaches include compartmental models [21, 55], agent-based and network models [36, 72, 78, 87], and curve-fitting models [11]. Additionally, **artificial intelligence (AI)** solutions, which have been successfully applied to complex tasks [71] such as image recognition [54, 124], speech recognition [14, 57], natural language processing [31, 58] and time series prediction [31, 58], have shown exceptional results in predicting future time series of observed cases and deaths much more accurately than classic regressive and simulation-based models (as documented in recent surveys [92]). However, despite better performance over more traditional modeling approaches, AI solutions are rarely adopted by decision-makers to support evidence-based policy-making and policy interventions related to disease outbreaks.

The limitations of AI for infectious disease mitigation were discussed at the panel on “The Future of AI for Spatiotemporal Data Science” held at the DeepSpatial’20 workshop co-located with the ACM KDD 2020 conference. Panelists, including three U.S. National Science Foundation Program Directors, discussed why AI solutions cannot be used to curb the COVID-19 pandemic. Given their data-driven nature, a key limitation of AI models is the lack of high-quality infectious disease data, especially near the onset of an outbreak, on which to learn and build AI models. Furthermore, it has been shown that inherent biases in model input data can be unintentionally propagated through the AI prediction pipeline [67], meaning that highly accurate model outputs may only reflect the model’s ability to predict trends in the data rather than true cases and deaths [67]. Given that input data for such models tend to be biased across geographic and demographic dimensions, such models may fail to generalize to populations that are underrepresented in the data the models are trained on [50].

For example, test data used as inputs for AI models that measure the prevalence of infectious diseases over space and time are subject to a variety of biases that are difficult to estimate and correct. For the case of COVID-19 case data, some of this bias stems from the willingness, access, or ability of certain groups to participate in testing. Participation in testing is influenced by symptom severity [50], symptom recognition [18], occupation such as healthcare workers [112], ethnicity [34, 77], frailty (susceptibility of more significant adverse effects) [56], place of residence [39], social connectedness [69], internet access [12], and medical/scientific interest [113].

Therefore, our vision is to use simulated worlds as a virtual testbed to better understand and minimize bias in AI models of infectious disease prediction to improve their reliability and fairness as a tool for decision support. Toward this vision, we seek to answer two questions:

- (1) For different infectious disease spread prediction models, can we measure and quantify the link between data bias and prediction bias? In other words, can we measure how susceptible (or robust) a given infectious disease spread prediction model is to different sources of data bias? For example, if we know that a geographic region is underrepresented in the data, then can we understand to what degree the predicted number of cases will be systematically underrepresented in the predictions?
- (2) If we know (or at least, can estimate) the data bias, then can we use the above understanding to account for and correct the systematic bias?

Given the explosion of AI solutions across a wide range of applications, this vision supports the increasing demand for transparency, fairness, and inclusion in AI and calls for increased consideration of bias propagation by the modeling community as a whole.

1.1 Stylized Example: How Data Bias Impacts Model Bias

To better illustrate the underlying problem motivating our vision, consider the following example that we will refer back to throughout the article:

Assume an outbreak of a novel infectious disease in a city with four neighborhoods. Testing kits to test, detect, and report the disease are available at grocery stores but at a high monetary price. **Ground Truth (Hidden):** In reality, at T_0 there are outbreaks in three neighborhoods of the city—one affluent high-income neighborhood and two low-income neighborhoods.

Observed Data: Using the testing kit results, we observe a large number of cases in the affluent neighborhood. Due to the high price of testing kits, we observe a small number of cases in the two affected low-income areas, and zero cases in the fourth neighborhood.

Two prediction models are developed to predict the spread of the disease (assuming no interventions) over time, as follows.

Prediction Model 1: A spatial compartmental **Susceptible-Infectious-Recovered (SIR)** model [61], which extends the traditional SIR so that the number of I in neighborhood i at time $t + 1$, is also a function of the number of I in adjacent neighborhoods j at time t and the transmission probability β between those in neighborhoods i and j . The model parameters β and the recovery rate γ are informed by domain expertise about infectiousness and transmission pathways. This model forecasts a uniform spread across all neighborhoods in two weeks.

Prediction Model 2: A predictive AI model (such as surveyed in Reference [92]) that aims to minimize the difference between predicted cases and the cases that will be detected in the next two weeks. This model is trained on the observed data and predicts a steady number of cases in affluent neighborhoods but a low number of cases in all low-income neighborhoods.

Two weeks later, at time T_1 the infectious disease has spread.

Ground Truth (Hidden): The infectious disease has spread uniformly across all neighborhoods.

Observed Data: High-income areas have a large number of observed cases. Low-income areas have a small number of observed cases.

The forecasts for Prediction Models 1 and 2 for Time T_1 are compared with the observed data. We observe that Prediction Model 1 has a large error: It predicted too many cases in low-income neighborhoods, compared to the observed data. Prediction Model 2, however, has a very low error. It successfully predicted the trends in the observed data.

This example illustrates the idea that data-driven approaches, such as Prediction Model 2, may demonstrate high prediction accuracy based on their ability to predict trends in the data. This is particularly problematic in the case of predicting infectious diseases, where the data used to train such models is highly biased [50]. Without understanding and correcting for biases in such data, we hypothesize that AI models may be less reliable for predicting the real trends in cases and deaths.

We acknowledge that Example 1.1 is highly stylized, making potentially unrealistic assumptions about bias in infectious disease data. The question is: How do different types of data bias found in real-world data propagate through the AI modeling pipeline affecting the prediction results? This is a particularly challenging problem, given that bias in real data is relatively difficult to quantify and isolate—a classic problem of “we don’t know what we don’t know.”

Therefore, to meet our vision of using simulated worlds for better understanding and minimizing bias in AI models of infectious disease prediction, we propose the following solution. First, an agent-based model of infectious disease spread that allows for *in silico*¹ collection of simulation data (both simulation ground-truth and sampled simulation data that is biased in some way). This simulation allows for the investigation of which AI models are robust to data bias and which models overfit to data bias. We will employ various scenarios of data bias across geographic and demographic dimensions to quantify how such biases impact AI model outputs, and thus we emphasize that we do not necessarily require that the sampled data is a perfect match to real biased datasets. Where not possible using real datasets, this approach allows us to have both a perfect measurement of the bias found in the sampled data and to isolate different types of bias.

Second, once the link between data bias and AI model bias is understood, we aim to combine (1) the understanding obtained about the link between data bias and model bias obtained *in silico*, and (2) knowledge about the bias in real data, to correct the bias in AI models for infectious disease spread. While this bias correction can also be performed in our *in silico* simulation, it will remain an open question whether the correction techniques learned in our simulated world also apply to the real world. *However, we hypothesize that the bias correction learned from our simulated world(s) may be better than simply ignoring the data bias and treating infectious disease samples as representative and unbiased.*

1.2 Vision of Simulating Bias in Infectious Disease Case Data

Given the challenge of estimating the type and magnitude of such bias in the variety of data sets used as model inputs, it is difficult to evaluate how this bias propagates through the AI model pipeline and how to mitigate this effect. Therefore, we envision using a massive agent-based model to simulate a (virtual) world in which we can collect perfect data with a 100% sampling rate and 100% accuracy of all information of the simulated world. We note that this simulated world may not reflect the real world, as “all models are wrong, but some are useful” [19] and may have their own biases [62]. We propose to use a realistic agent-based simulation based on realistic patterns of life such as going to work, going to restaurants, and meeting friends to socialize. Informed by real-world data about the environment (road network and buildings of a city), we hope that such a simulation, while not being a true representation of the real world, may be representative of some aspects of a real population.

Based on these data, we can then implement sampling strategies with different types of bias common to both the sampling methods used to collect these data and the participation of simulated individuals (agents). Using the example of Section 1.1, the simulated world allows to increase

¹An *in silico* experiment refers to an experiment performed in a computer simulation and is an alternative to experiments involved *in vivo* (on living beings) and *in vitro* (in a laboratory).

(decrease) the participation of high-income (low-income) individuals to produce different data sets with different levels of data bias. Applying state-of-the-art AI-based spatiotemporal infectious disease prediction to this sampled (and biased) data set will allow us to compare predicted results to the results obtained using the full (or unbiased) data set. While this bias that we simulate may or may not perfectly reflect the real-world, it still allows us to answer the research question of “If there is a certain type X of bias in the data, then the spatiotemporal infectious disease predictions become biased by Y .” Or formally:

$$[\text{Data Bias } X] \rightarrow [\text{Model Bias } Y].$$

Using our example in Section 1.1: Assume that we scale the bias at which data is collected from the high-income and low-income neighborhoods in our simulation. We evaluate Model 1 (the spatiotemporal compartmental SIR model) and Model 2 (the AI model) using this simulation by comparing prediction results (based on the biased simulation data) to the simulation ground truth. Assume that we observe the following implications:

$$\text{Model 1: } [\text{High Data Bias}] \rightarrow [\text{Low Prediction Bias}];$$

$$\text{Model 2: } [\text{High Data Bias}] \rightarrow [\text{High Prediction Bias}].$$

That is, we find that Model 1 is robust to the simulated type of data bias and still make accurate prediction; while Model 2 overfits to the simulated data bias and leads to strongly biased predictions.

This result, by itself, does not allow us to make any inference on the real world, where we do not know if the simulated data bias holds. The bias that we simulated is an assumption that may or may not hold in the real world. Implications ($A \rightarrow B$) such as the above (if inferred in the *in silico* simulation) still hold even if the premise (A) does not hold (logically: *ex falso quodlibet*). But if the premise (A) does not hold in the real world, then the implication ($A \rightarrow B$) does not yield any information, as both $A \rightarrow B$ and $A \rightarrow \neg B$ hold logically if A is false.

But now, assume that for a specific infectious disease data set observed in the real world, we know that the observed case data is highly biased toward high income. In other words, we know that the premise (A) holds. Then, the implication ($A \rightarrow B$) allows us to infer the conclusion (B). In the example, if we know that a given dataset is highly biased toward a high income (premise [High Data Bias]), then we can use the two implications above to infer that:

$$\begin{aligned} \text{Model 1: } [\text{High Data Bias}] &\rightarrow [\text{Low Prediction Bias}] \wedge [\text{High Data Bias}] \\ &\rightarrow \text{Model 1: } [\text{Low Prediction Bias}]; \end{aligned}$$

$$\begin{aligned} \text{Model 2: } [\text{High Data Bias}] &\rightarrow [\text{High Prediction Bias}] \wedge [\text{High Data Bias}] \\ &\rightarrow \text{Model 2: } [\text{High Prediction Bias}]. \end{aligned}$$

This knowledge will allow decision-makers to understand which models to trust given their knowledge about real-world data bias and the *in silico* inferred link between data bias and model bias.

1.3 Overview of the Proposed Vision

Figure 1 presents an overview of our proposed vision for a novel methodology that can assess and mitigate bias in AI solutions. The methodology is made up of two components, as follows:

- (1) **Developing a scalable agent-based model from which we will collect simulation data with controlled bias.** Our vision is to leverage scalable agent-based simulation [7, 63, 85, 87] to create a simulated sandbox world from which data sets are collected under various bias scenarios. Such simulation will be able to model a real region (such as Fairfax

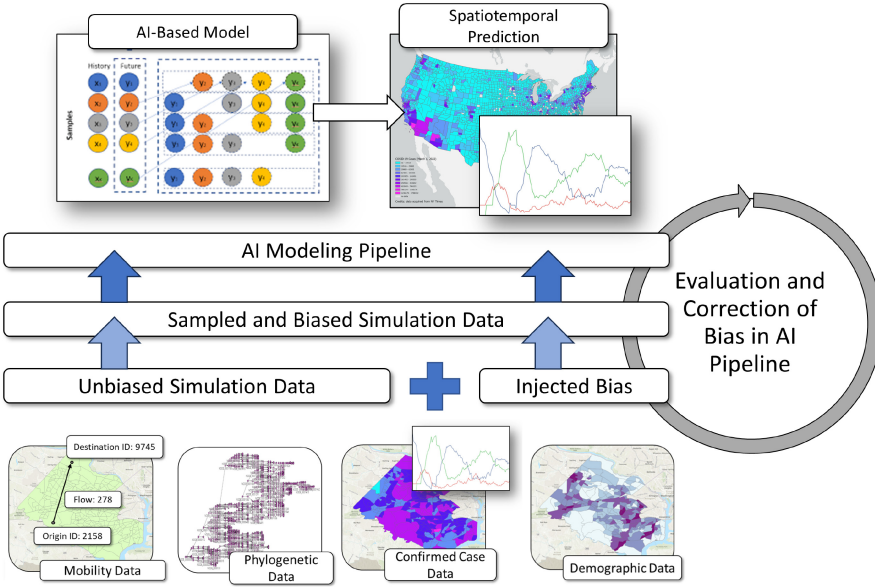


Fig. 1. Overview of the vision.

County [87] or New Orleans [63]) using large sets of human mobility data and census data to create a realistic digital twin of the region. Within this world, we can leverage state-of-the-art privacy solutions [75] to emulate privacy-aware data collection or aggregation to obtain realistic data sets from the simulation. Informed by public health experts, we will use the simulation to (1) simulate the spread of an emerging aerosol or sexually transmitted disease, and (2) simulate various forms of data collection bias.

- (2) **Understanding and correcting bias in AI systems.** The simulation sandbox will allow us to explore links between bias in data collection and the resulting bias in an AI model prediction based on such biased data, and investigate solutions to correct such bias. Leveraging recent work on fairness in AI [102, 104], a sandbox world will allow us to evaluate existing solutions for measuring bias and fairness in predictive models induced by different types and levels of bias in the simulated data collection. As one source of bias is privacy mechanisms during data collection and aggregation, the sandbox will allow investigating solutions to mitigate the data bias induced by different privacy constraints. In addition, we can also investigate privacy-enhanced predictive models using the raw simulated data (as versus adding privacy at the data collection or aggregation stage) and how they may amplify the bias in the data. Once bias and fairness in data are understood, we can leverage work on class-contrastive counterfactual explanations [88, 103, 105] to explain and mitigate bias.

In the following section, we provide a brief overview of different state-of-the-art models for data-driven infectious disease spread prediction. Section 3 provides further details on our vision of using a sandbox agent-based simulation to simulate and control bias in data collection. Section 4 presents a prototype simulation and benchmark bias dataset. Section 5 details the vision of using this sandbox to understand and mitigate the prediction bias that may be amplified from the simulated data bias for different prediction models. Broader potential societal impacts of our vision are described in Section 6.

2 RELATED WORK ON UNDERSTANDING BIAS IN INFECTIOUS DISEASE SPREAD PREDICTION MODELS

In this section, we survey state-of-the-art methods for data-driven infectious disease spread prediction (Section 2.1) as well as methods for handling bias in infectious disease models (Section 2.2)

2.1 Existing Models for Data-driven Infectious Disease Spread Prediction

Data-driven epidemic forecasting has been a very large research field in the past decade. A recent survey summarizes more than 300 publications in this field [92]. Even before the COVID-19 pandemic, this field was already a focus of the computing community [79] in the context of influenza-like illnesses [2]. For example, ACM KDD has been organizing the International Workshop on Epidemiology meets Data Mining and Knowledge Discovery since 2018 [1, 3]. The COVID-19 pandemic has brought forth very large sets of human mobility data [39, 45, 90], which enabled new data-driven models. The ACM SIGSPATIAL has been organizing the International Workshop on Modeling and Understanding the Spread of COVID-19 [10], the International Workshop on Spatial Computing for Epidemiology [8, 9], and community papers on challenges for Mobility Data Science [83, 84], which include improving the understanding of infectious disease spread using mobility data. Existing data-driven models to predict the spread of infectious diseases include compartmental models [4, 61], agent-based models [87, 114], regression models [47], off-the-shelf sequential models [115], graph neural network models [32, 119], density estimation models [20], ensemble models [26, 91], contrastive predictive coding [108], as well as many other types of models [92]. Despite the plethora of data-driven models to predict the spread of infectious diseases, our understanding of how bias in data collection may affect the accuracy and reliability of resulting predictions of different models is lacking.

2.2 Existing Approaches to Handling Data Bias in Infectious Disease Models

Studies that focus solely on minimizing data bias to support analysis and modeling are numerous [50]. There are, however, few studies that explicitly describe their approach to handling data bias in their models of infectious disease spread. Take traditional compartmental models for example; most compartmental models assume perfect reporting of infected cases, resulting in incorrect model parameters when fitting the model to the data [44]. Techniques such as Bayesian inference [106], maximum likelihood estimation, and data smoothing and interpolation [74] can help identify and adjust for reporting biases in the data. Additionally, a few studies investigate the impact of bias using simulated data on parameter estimation and prediction using compartmental models. For example, Suhail et al. [107] find that bias toward testing symptomatic individuals increases the predicted number of cases over time. Krishnan et al. [68] investigates the impact of different spatial aggregations of case data on R_0 estimation, which is typically an important parameter in models to gain insight into the early growth rate of an epidemic.

Recently more attention has been given to the problem of the natural changes within populations that occur as a consequence of the epidemic itself, and how traditional infectious disease modeling approaches largely ignore or are incapable of modeling such interactions. Meadows et al. [82] reviewed and analyzed the issue of input data bias and its impact on disease modeling during the COVID-19 pandemic, and highlighted the importance of correction of biases prior to model parameterization. Equally, evolving biases such as behavioral changes in human populations (self initiated reductions in mobility, reduced contact, social distancing) as they become aware of a pandemic over time can lead to substantial discrepancies between estimates from models and real-world epidemiology for the estimation of key epidemic parameters such as R_0 [38]. Correcting for evolving biases within models is most easily dealt with in individual-based models, while

deterministic and stochastic models can only address such effects using larger-scale population approximations.

To handle bias in AI models, researchers use different approaches depending on the source and type of bias, the goal and context of the application, and the ethical and legal implications of the bias. Some of the common approaches surveyed in References [97, 101, 123] are: (1) Pre-processing the data to correct for known bias, such as up-sampling or downsampling datasets. Such an approach is challenging in the context of infectious disease spread where the bias (e.g., the testing rates of different neighborhoods) is not known *a priori*; (2) Post-processing the output of models to adjust for known data bias, which is also difficult when the underlying data bias is not known, and; (3) incorporating fairness constraints by rewarding models to yield a fair result. In the context of infectious disease spread, this may also be difficult or even dangerous, where incurring deliberate false predictions for fairness may mislead public health decision-makers.

In general, tackling bias in prediction models appears to be largely overlooked in the literature and typically lacks a focus on social fairness. Additionally, despite such important efforts at correcting for data bias, the multifaceted and complex nature of such bias means that it is impossible to perfectly quantify so that its effects in the modeling pipeline can be fully studied. The following describes our vision of using a sandbox world in which we can (1) control the bias in observable data, (2) evaluate the robustness of different models to data bias, and (3) investigate solutions to mitigate prediction bias.

3 AGENT-BASED MODEL FOR SIMULATING DATA WITH CONTROLLED BIAS

To measure and correct the effects of bias in the AI modeling pipeline, we require a perfect understanding of the bias inherent to data commonly used as input for AI models. Ideally, we would need a complete understanding of the real system (i.e., infectious disease spreading through a set of individuals over space and time) and the degree to which the data that are captured from the system are representative. Therefore, the objective is to build a scalable **agent-based microsimulation (ABM)** that will serve as a sandbox world of which we have perfect knowledge and can collect data that capture characteristics, mobility, and disease prevalence of 100% of individuals—the Ground-truth data. In addition, we can collect samples of such data under various bias scenarios, which we refer to as Observable Data. In Figure 2, we show an overview of our methodology of simulating both the sandbox world and the data collection within this world. While real data sets have unknown biases that are challenging to measure, ABM provides the opportunity for different types of bias in the Observable Data to be injected and measured so that its propagation through the AI modeling pipeline can be thoroughly investigated. It should be noted that the purpose of ABM is not to be a perfect representation of the real world, but rather a complete representation of a system from which data can be collected. For this purpose, existing solutions for scalable epidemic simulations can be leveraged [7, 63, 64, 85, 87, 125].

We can extend existing ABMs of disease spread [87] that are both flexible and general so that they can be re-parameterized to simulate a variety of infectious human diseases that are transmitted through droplet spread or direct contact (e.g., sexually transmitted diseases). We can utilize a synthetic population generation approach to generate synthetic individuals and their characteristics both in terms of socio-demographic make-up (such as gender, race, age, sexual orientation, and the like) and daily activity sequences. We can modify the selected agent characteristics and activities as needed based on the model application (i.e., the disease that is modeled, and the research questions). Based on the synthetic population's characteristics, behaviors, and mobility, physical contact and sexual contact networks can be estimated along which the diseases may be transmitted.

We can simulate the collection of four data types, derived from the synthetic population in the ABM: (1) mobility data, (2) census data capturing socio-economic profiles of various regions,

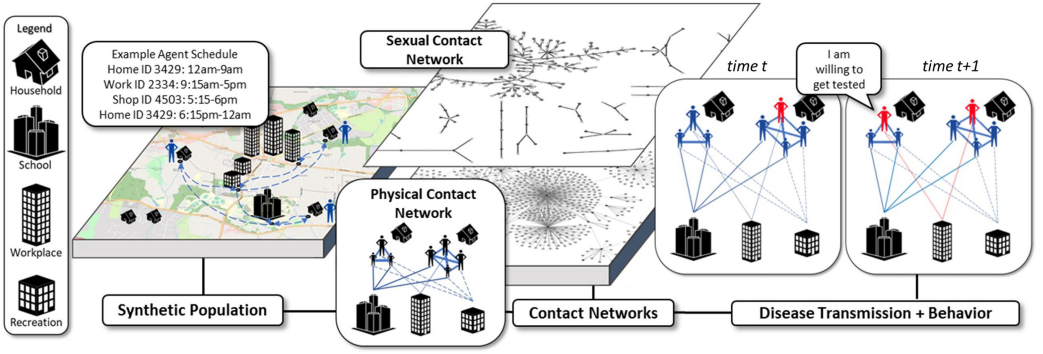


Fig. 2. Overview of agent-based modeling framework.

(3) confirmed case data, and (4) phylogenetic tree data. These data will be used as the input for existing AI models to test and correct for bias. For each observable data set, we can first simulate the sampling bias inherent to the data. Each of these open data sets can be considered volunteered information, meaning that synthetic individuals in the ABM decide whether or not to own a smart-phone and share their location data, whether to participate in the census, whether to get tested, and whether to report symptoms. Sampling bias can occur in this type of data when the groups of individuals that choose to volunteer their information differ from the individuals that do not and thus the data are not representative of the entire population. This can emerge in the case where certain individuals lack trust or access to the systems that collect the data.

The simulation will allow for sampling bias scenarios where the populations that volunteer their information can be modified through a variety of hypothetical scenarios or estimated based on existing studies. In addition, we can apply privacy mechanisms including data aggregation and noise addition, calibrated to meet the standards of **differential privacy (DP)** [28, 37]. Privacy mechanisms can unintentionally create additional forms of bias in data. Recent work has demonstrated that when privacy-protected data are used for downstream decision-making as if they were true, different population groups may be treated differently [89]. Other works have shown that privacy protection can mask statistical disparities and thus conceal evidence of the disparate impact that is potentially discriminatory [122]. Privacy-preserving models also have a disparate impact on model accuracy, i.e., larger misclassification rates for underrepresented groups compared to well-represented groups [15].

Expertise in spatiotemporal data privacy [23, 24, 51–53, 111, 116, 120] can help inform the data collection and aggregation process to apply various differential privacy mechanisms in the generated data. For example, data on cases may not be published (and thus, observable) for spatial regions (such as census blocks in the United States) having a population below a certain threshold to avoid identification of individuals; or different level of noise may be added to different regions depending on the case frequency. Such data privacy mechanisms and policies can be implemented in the simulation to gain an unprecedented insight (as in the real world, we do not know what data we do not have) as to how such simple privacy policies may incur data bias and how this data bias may affect infectious disease models. We can also simulate more complex privacy-preservation approaches such as differential privacy to create observable data sets that have a similar bias as observed in the real-world described in the following paragraphs.

Human Mobility Data. Human mobility data capture the movement of individuals from a set of origins to a set of destinations and are commonly used as input for models that simulate the spread

of disease across geographic space and time [87]. Individual-level mobility is captured by mobile phone data, collected actively through call detail records and passively through smartphone applications. These data are typically anonymized and aggregated to produce a range of mobility data products at various spatial resolutions [13, 33, 43, 45, 49, 96]. For example, SafeGraph [96] estimates the number of individuals that move from a census block group to a set of commercial points of interest. Publicly available mobility data sets capture a small fraction of the population (less than 5%) and thus tend to be subject to *sampling bias*. For example, the SafeGraph data are biased toward populations that use smartphones. In addition, data may be subject to noise perturbation for privacy protection. For example, the SafeGraph data set applies differential privacy approaches to home census block groups of individuals as well as to visited locations [95].

Demographic Data (Census Data). The U.S. Census Bureau's **American Community Survey (ACS)** captures the demographic makeup of populations at various geographies ranging from very large at the state level to very small at the census block group level. In disease response and mitigation, these data are commonly used to identify vulnerable populations that may be more susceptible to outbreaks or may need tailored interventions. These data can also be used to measure associations between population demographics and cases, deaths, vaccine uptake, and so on. ACS data are subject to *response bias* where populations who complete the census differ from those who do not, referred to as non-respondents. Non-respondents are typically greater in areas with non-White populations, and areas having lower household incomes, less home ownership and fewer college graduates, thus these populations may be underrepresented in census data [94]. Typically, demographic information becomes less detailed as resolution increases to preserve privacy. Individual level data are only available in the Public Use Microdata Sample, which is a 5% sample of the population within a large Public Use Microdata Area.

Confirmed Case Data. Health departments and other agencies collect counts of confirmed cases and may make these data available, aggregated at various spatial resolutions. There are many sources of bias in the collection of epidemic data in public health [40, 118]. One is testing bias, where some countries have better testing infrastructure, well-funded access to testing and less stigma around getting tested. Such countries may appear to have more cases than a country where a lack of resources or stigma results in low testing rates [81]. Another source of bias are case definitions. It is reported from many countries that some population groups, such as women and children, have greater barriers to accessing testing or vaccination [93, 109]. This may bias the case data toward men, and more importantly, miss spillover transmission into wider population groups. The presence of asymptomatic infection may also lead to a substantial underestimation of the case count if only symptomatic people receive testing. In the case of SARS-CoV-2, up to 30% of cases have asymptomatic infections [98], leading to undetected community transmission.

Phylogenetic Data. Originally an initiative to foster the collaborative sharing of influenza virus data, **Global Initiative on Sharing Avian Influenza Data (GISAID)** is a trusted and widely used platform for the rapid sharing of sequence data for a variety of diseases such as influenza, monkeypox and SARS-CoV-2. The database captures virus sequences as well as related clinical, epidemiological and geographical metadata [100]. Users can then download both sequence data and metadata to develop phylogenetic (evolutionary) trees. In particular, users may generate a spatially and temporally resolved tree that shows the evolutionary diffusion (spread) of sampled viruses to their **most recent common ancestor (MRCA)** [16]. In the case of a pandemic, the estimated location of the MRCA may suggest the origin of then outbreak [73]. We note that there is obvious sampling bias in the available sequence data. For example, during the SARS-CoV-2 pandemic, the United Kingdom and the United States at one point represented 61% (38% and 23%, respectively) of

the SARS-CoV-2 sequences in GISAID [100], whereas after Denmark (6.6%), every other country contributed less than 4%. Based on the recorded case data, this means that GISAID has, on the surface, adequate representation of the U.S. cases (since it has 25% of the global cases), but inadequate representation of countries like India and Brazil, each of which has about 10% of the total worldwide amount but well less than 1% of the sequences in GISAID. To account for sampling bias, we will proportionally sample GenBank and GISAID at the country level based on known case counts [35]. In our sandbox, we will have a complete phylogenetic tree. However, the data set that is produced from our sandbox will be a subset of case data and an even further subset of cases that are sequenced.

Existing algorithms for infectious disease spread forecasting will then be applied to both Ground-truth and Observable Data as described here to understand how biased data lead to biased predictions. We can then evaluate bias for different regions and populations, in particular populations underrepresented in the data.

4 PROTOTYPE SIMULATION AND BENCHMARK BIAS DATASET

Acknowledging that this is a vision article, we provide a preliminary prototype study to show (1) how an *in silico* simulation can be used to synthetically scale data bias, and (2) to provide a simple benchmark dataset to scale the bias at which data is collected. This scenario used in this prototype study follows the running example described in Section 1.1. In this section, we first survey an existing agent-based simulation used for this prototype in Section 4.1. Then, Section 4.2 describes how this simulation was applied to the city of Atlanta, Georgia, USA for the prototype, how a simple disease model was imposed on the simulation, and how data collection bias for a high-income area was imposed on the simulation. Section 4.3 describes the generated datasets and evaluates the imposed bias.

4.1 Patterns of Life Simulation

Our prototype is based on a socially plausible agent-based simulation called Urban Life [125]. Urban Life is an agent-based city-level simulation in which each agent represents a simulated human in the real world that follows socially plausible patterns of life. The simulation allows to leverage real-city environment data (road network, buildings, apartments) leveraging a pipeline to extract data from **OpenStreetMap (OSM)**. Details how the simulation can be adopted for any region in the world using OSM data can be found in Reference [6]. Agents in this simulation correspond to individual humans who commute between their home and work locations. Agents go to restaurants to eat and go to recreational sites to meet friends and socialize. A social network that captures friendship, family, and co-worker relationships evolves as agents interact with each other over time.

Agent behavior is driven by Maslowian needs [80] such as physiological needs (shelter, food), financial needs (money), and love needs (friends, family). These needs drive the decision-making of agents that lead to behavior to satisfy the needs, leading to an emerging behavior in which agents find a balance between spending time and making money, meeting friends, and satisfying other needs. An in-depth description of the Urban Life simulation can be found in Reference [125] and the Java-based source code of the simulation can be found on GitHub at <https://github.com/gmugs/pol>.

4.2 Infectious Disease Model in the Patterns of Life Simulation

To simulate the outbreak of an infectious respiratory disease we augment the Urban Life simulation with a simple SIR disease model. Initially, all agents are in the susceptible state. A small number of agents is initially set to infectious. Agents that are susceptible become infectious through

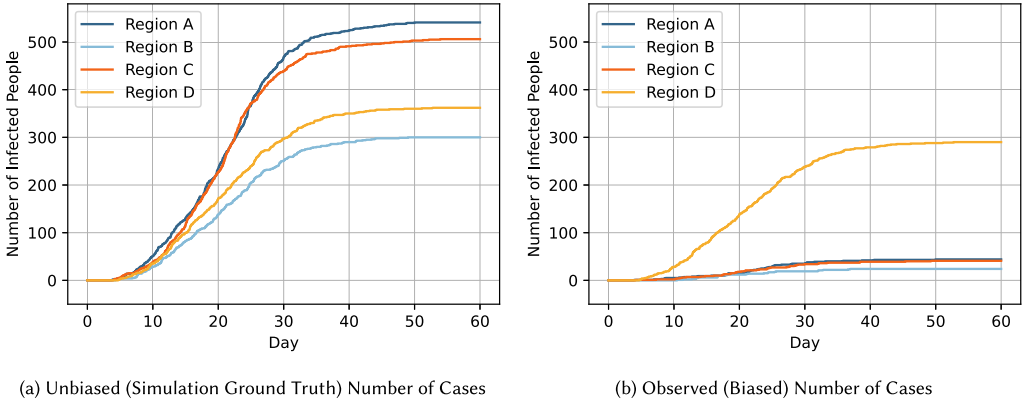


Fig. 3. Example dataset of a simulated infectious disease outbreak in Atlanta, Georgia, USA. Different datasets with varying degrees of sampling bias can be found at <https://github.com/RuochenKong/disease-simulator>.

exposure to another infectious agent (see details below). Infectious agents will automatically become recovered after seven simulation days. Recovered is a terminal state.

To simulate exposure between agents, the Urban Life uses the concept of meetings in which two or more agents interact in person to increase their social ties. Meetings operate as the main transmission pathway, allowing for the spread of an infectious respiratory disease. In general, respiratory diseases require a longer duration of close-contact exposure. Therefore, in the simulation, agents must be in a meeting for at least 5 min to become exposed. Thus, every 5 min in a meeting, any infectious agent in the meeting has a 1% chance of infecting any susceptible agent in the same meeting. In the simulation, meetings mainly happen at home and recreational sites.

We note that in this prototype simulation, agents do not change their behavior once they are infectious. In particular, agents do not use any mitigative actions or avoid meetings when they are infected. We note that such behavior change can be implemented in the Urban Life simulation as described in Reference [66]. But for this prototype, we want to keep the data generation simplistic.

We have made the resulting simulation, including the infectious disease model, in a GitHub repository found at <https://github.com/RuochenKong/disease-simulator>.

4.3 Prototype Biased Data Benchmark

To obtain biased case data from the Urban Life simulation described in Section 4.2, we follow the narrative of the running example described in Section 1.1 whereby one spatial region is sampled at a higher rate than other regions. We divide the Atlanta study region into four approximately equally sized (in terms of population) regions. We define one of the four regions (Region D) as having a higher sampling rate. In the running example, this corresponds to the high-income neighborhood. We run the Urban Life simulation (Section 4.1) having with 2,000 agents. Figure 3 shows the resulting number of cases, both unbiased (Figure 3(a)) and biased having a sampling rate of 80% in Region D and a sampling rate of 8% in Regions A, B, and C.

This dataset, as well as other datasets with different sampling rates available at our Github (<https://github.com/RuochenKong/disease-simulator>) can then be used as a benchmark for different prediction models to experimentally measure the function that maps data bias to prediction bias. This function can not be observed in the real world allowing this simulation-based approach to better understand the robustness of different prediction models to the underlying data bias.

5 UNDERSTANDING, MEASURING, MITIGATING, AND CORRECTING BIAS

The simulation framework will provide us with a sandbox world that we can use to generate and collect data to feed to pre-existing AI models (surveyed in Section 2). The advantage of employing the aforementioned sandbox over using real-world data is our ability to control the data collection strategy in addition to having unrestricted access to an authoritative ground truth of the full census data that spans the entire population. Such a setup allows us to evaluate how injecting different types and strengths of bias into the data generation and collection procedure implemented within our simulated world affects the AI models in terms of (re)training, tuning and predictive ability. Additionally, full control over the modeling pipeline provides us with a unique environment to investigate the suitability of diverse methods to identify, measure, mitigate and correct bias across a range of distinct scenarios. Specifically, we want to answer the following questions:

- How do different types of human mobility, case report and phylogenetic data bias affect the ability of AI models to learn fair, robust and predictive representations?
- How can we correct the bias in AI models for a selection of known biases introduced into the data generation and collection process?
- How do these biases and their correction mechanisms impact the predictive performance of AI models overall and across individual populations?
- How can we better estimate the true error of AI models without knowing the (source of) bias in our data?

The major advantage of using a sandbox environment is the ability to repeat the simulation with different starting conditions and parameters to obtain counterfactual worlds. This allows us to implement different patterns of bias in the data and evaluate how these conditions affect the predictive performance and bias of the AI models when used as a (re)training or tuning sample. By repeating experiments a sufficient number of times we can gain insights of statistical significance, which is of paramount importance for a high-stakes setting such as this one.

5.1 Understanding the Links between Data Bias and Modeling Bias

First, we want to understand how well-defined (and controlled) data bias affects the bias of the AI models when (re)trained or tuned on these data. We will inject bias into data sets iteratively over many simulation runs, magnifying a single bias source and mixing together different types of bias. Each simulation will have data sampling bias parameters, such as the degree to which certain populations—stratified by age, gender, race and income—are sampled in terms of human mobility, disease reports and viral genome sequencing. Depending on the task, each simulation output will then be fed to the infectious disease spread predictive AI models both to (re)train or tune them and to provide us with predicted cases. By repeatedly running our simulation with different data bias parameters we can obtain a database linking *simulation parameters*, *data*, (ground-truth) *labels* and *predictions*. This database can be mined for patterns and correlations; in particular, we can evaluate the following hypotheses:

Spatial Aggregation Undersampling a spatial region—for example, caused by a low population in that region leading to data gaps due to location privacy—may result in the number of cases being underestimated in that region, as the underlying AI model is unaware of the population due to sparsity of training data. In addition, the choice of spatial partitioning of the universe of discourse into regions—a problem known as the Modifiable Areal Unit Problem [42, 46, 117]—may affect the AI model. Some attempts have been made to address this, although more work in this area is needed [121].

Socio-demographic Population Undersampling a socio-demographic population—such as the group aged 65+—may lead to an underestimation of cases in regions having a large

representation of this population. Alternatively, there might be an overestimation due to the AI model (incorrectly) extrapolating cases from a smaller population, thus increasing the sample variance.

Overrepresentation Having a large sample size of cases in a region—for example, due to free testing being offered in a given location or for a particular population—may lead to having more cases of an infectious disease in the observed data, which does not necessarily reflect the true distribution, possibly creating an overestimation of spread in that region.

Data Quality Data collection is often subject to errors, noise and missing information with respect to both instances and labels. Poor quality of data—either throughout or for a particular subpopulation—may cause the model to underperform or learn a biased representation, which will be reflected in the number of estimated cases.

Given the flexibility and full control over data generation, collection and modeling we can easily investigate these scenarios and assess their severity and impact. Specifically, we can compare data generated without any bias with samples polluted by a range of known biases that are processed with dedicated correction algorithms. The same procedure may be followed when dealing with state-of-the-art AI models, allowing us to test a wide array of pre-processing, in-processing, and post-processing methods [5, 17, 22, 27, 41, 59, 65, 70, 76, 102]. Such a test bed allows us to perform a comprehensive study and evaluation of available techniques, and offers a development environment for novel methods. Among others, our work will cover four mainstream notions of group fairness [25]—*demographic parity*, *equalised odds*, *false positive* and *true positive*—paired with three types of bias correction algorithms based on: pre-processing [17, 22, 41], in-processing [5, 59, 70], and post-processing [27, 65, 76].

5.2 Robust Solutions for Fairness in AI Models

Once we understand the sources of bias and identify methods that are appropriate to measure it, we can investigate bias correction mechanisms suitable for various bias scenarios, and their robustness to controlled changes in the data distribution and quality. For example, we can investigate how the detrimental side effects of (pre-)processing techniques common to this type of (sensitive) data, e.g., intended to preserve privacy of individuals, can be effectively counteracted with bias correction algorithms. We may also look into augmenting real-life data or supplementing them with simulated samples by studying mixtures of samples generated with different initial conditions to identify methods capable of improving the robustness of state-of-the-art data-driven AI models for infectious diseases spread prediction. In addition, we can investigate adding state-of-the-art privacy-preserving machine learning techniques [28] for disease predictive models and how it may amplify the bias in the data. We may also revisit the privacy mechanisms and investigate if we can develop privacy mechanisms that are equitable or satisfy the fairness notions we discussed above. It has been recently shown that there are cases where the constraint of differential privacy precludes exact statistical notion of fairness (namely, equality of false positives and equality of false negatives) [29]. However, it is possible to build a classification algorithm that maintains utility and satisfies both DP and approximate fairness with high probability. We may investigate data aggregation and learning algorithms with both DP and fairness constraints with approximations in either or both of them as appropriate, and then evaluate their impact and how they compose with the other bias correction algorithms we will develop.

5.3 Evaluation and Error Estimation of Corrected Models

To achieve our vision, it is essential to assess the impacts of biases and correction mechanisms on AI models, as well as to estimate their performance beyond simulated environments without

prior knowledge of real-world data sources and distributions. Analyzing the impacts of biases and correction mechanisms on AI models across various scenarios is important in understanding how different types and strengths of biases could affect the performance and fairness of the AI models. Leveraging the simulations, we could conduct sensitivity analysis, which enables the evaluation of the robustness of corrected models to changes in data distribution and quality. Varying input parameters and data in the simulation environments allow for the assessment of model performance under different conditions and help identify potential vulnerabilities or areas for improvement. For comparing model outcomes before and after correction, metrics such as disparate impact and accuracy across different demographic groups could quantify the effectiveness of bias correction mechanisms. To be more specific, by conducting this analysis, we could potentially identify the most effective correction mechanism for a particular bias scenario. Further, based on these results, Mixture-of-Experts-like architectures would be a valuable and interesting working direction to enhance the AI models in terms of mitigating the potential unknown biases in real-world scenarios. Such sensitivity analysis will also provide valuable insights into the resilience of corrected AI models and guide further debiasing developments.

Estimating the error of AI models (including the corrected AI models) without knowing the bias in real-world data is a challenging task due to the hidden unexplored bias, complex interactions beyond simulation, and limited and sparse data sources. To address this ongoing challenge, several directions could be considered. First, robust data sampling and feature engineering: active sampling could be introduced to strategically select and prioritize data samples that have been largely investigated in the simulated environments for model training and evaluation. By focusing on informative data points, active sampling can help mitigate the effects of hidden biases and sparse data sources. Besides, advanced feature engineering approaches including using existing powerful LLMs (that have general knowledge and common sense about the open world) to extract informative semantic features could reduce the potential biases caused by the straightforward spatiotemporal feature selection process. Second, continuous updating with expert feedback: implementing mechanisms for continuous updating of models based on expert feedback could be beneficial in reducing the bias caused by limited data sources and exploring newly occurred biases. Similar to leveraging RLHF [60] (Reinforcement learning from human feedback) used for improving LLMs, we could also design new correction mechanisms based on the feedback from experts. Additionally, we could reproduce the observed biases from the feedback in our simulation environments to generate more simulation data to provide high-quality data for investigating the correction mechanisms.

6 DISCUSSION AND CONCLUSION

Our vision focuses on investigating the way in which bias propagates through the data to models and predictions derived from AI-based approaches. Here, we discuss the primary contributions should this vision be achieved.

6.1 Rapid Response to New Emerging Diseases

There is potential to consider a range of infectious diseases including COVID-19, the 2022 global monkeypox outbreak, and H5N1. For example, the recent emergence of monkeypox offers a robust case study and motivation to investigate this problem and address questions of fairness, given that mpox has been longstanding in African countries, despite a surge in incidence since 2017 [86]. The spread of mpox in non-endemic countries in 2022 is unprecedented, complex and has a different pattern to spread in the African continent [110]. As of February 2023, GISAID captures 5,137 sequences of hMPXV—the human mpox virus [48]—which provide a solid sample for contrastive learning of AI models. However, since monkeypox is transmitted through close and often intimate

contact with most cases reported in gay, bisexual, and other men who have sex with men [30, 110], we will need to investigate different approaches to simulate data collection bias.

6.2 Improving Understanding of Bias in Disease Spread Prediction Models

We will develop and make available datasets that explicitly incorporate bias, such as overreporting COVID-19 tests in specific locations (e.g., Section 1.1) or underreporting in certain socio-demographic groups. Such datasets are critical to support the significant advancement of all disease spread prediction models, providing ground-truth data alongside biased samples, enabling researchers and modelers to directly investigate and quantify the impact of biases on their analyses and outputs in any data-driven model.

This approach is crucial for enhancing the awareness and understanding of bias within data-driven models, particularly in public health contexts. By exposing and quantifying the effects of bias, these datasets empower modelers to develop more robust, fair, and accurate predictive models. Furthermore, they offer a valuable educational resource for training the next generation of data scientists in recognizing and mitigating bias, ultimately leading to more informed and equitable decision-making processes in various sectors, including healthcare, policy-making, and beyond.

6.3 Trust in AI Models for Disease Spread Prediction

Beyond understanding bias in existing predictive models, we hope that our project will help decision-makers to improve their trust in AI models used for infectious disease spread predictions. By being able to quantify, evaluate and correct the bias incurred by such model, we hope the general trust in AI models for disease spread prediction will improve. Or at least, by quantifying sources of bias in AI models we hope to identify a path forward to improve AI models in future research to potentially enable the use of AI models for infectious diseases prediction by decision-makers in the future. This includes a greater understanding of the aspects of the model that are more robust and reliable, and thus, more trustworthy and safe to be automated and deployed, and the aspects in the decision making pipeline that require further scrutiny and human experts in the loop [99].

6.4 Understanding Region-Specific Data Bias

Our proposed approach will enable us to better understand how cultural, demographic and regional differences, may yield different biases in AI model predictions. Consider Example 1.1, where a large number of cases are observed in a high-income neighborhood. Without accurate information on the actual number of cases, it is challenging to determine whether the outbreak is confined to this affluent area (Option 1) or if it has also spread to low-income regions but remains unobserved due to barriers such as the high cost of testing (Option 2). In societies where there is minimal interaction between affluent and low-income groups, Option 1 might be more likely. However, in places like the United States, where people of varying incomes share common spaces like grocery stores and healthcare is costly, Option 2 could be more probable. By incorporating social science expertise into the project, we can simulate cultural differences, and thus investigate and correct different types of regional bias in our sampled data.

6.5 Enriching Curricula with Case-Studies on Fair AI

Programs across the world offer courses on ethical issues in their computer science undergraduate and graduate curricula. If our vision became reality, then we could enrich such courses through a demonstration framework that uses a (simplified) version of an agent-based simulation (such as envisioned in Section 3) that allows to change the bias at which individual-level data are collected and to evaluate how this data bias affects (simple and complex) predictive algorithms. For

example, the framework may allow users to decrease the sampling rate of the female population or increase the chance that a female agent will report symptoms when self-testing negative. The algorithm may show that such biased sampling may lead to an overestimation of the infection risk of female agents (and, symmetrically, an underestimation of the infection risk of male agents). The framework would implement simple statistical models (such as Bayesian classifiers), simple machine learning models (such as decision trees), and state-of-the-art deep learning approaches. By showing the effect of data collection bias on different models, students will be reminded to consider the assumptions made by algorithms (such as i.i.d. sampling), develop an intuition on how to interpret results when data collection bias is known or suspected, and understand the robustness against data bias of different models.

In conclusion, our vision is a novel approach that will allow for a better understanding and mitigation of bias in the AI disease modeling pipeline. Achieving this vision would mean improving trust in such models for informing fair and equitable policy interventions for rapid response to new emerging diseases, improving bias detection methodologies, and encouraging consideration of bias for more fair and equitable AI. Beyond this study, our vision calls to the broader data science and modeling community to investigate and mitigate effects of data bias. This is critical to ensure that data driven AI models are transparent, fair, and equitable tools for decision support and policy making.

REFERENCES

- [1] Bijaya Adhikari, Ajitesh Srivastava, Sen Pei, Sarah Kefayati, Rose Yu, Amulya Yadav, Alexander Rodríguez, Arvind Ramanathan, Anil Vullikanti, and B. Aditya Prakash. 2021. The 4th International Workshop on Epidemiology Meets Data Mining and Knowledge Discovery (epiDAMIK 4.0@ KDD2021). In *Proceedings of the 27th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 4104–4105.
- [2] Bijaya Adhikari, Xinfeng Xu, Naren Ramakrishnan, and B. Aditya Prakash. 2019. Epideep: Exploiting embeddings for epidemic forecasting. In *Proceedings of the 25th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. 577–586.
- [3] Bijaya Adhikari, Amulya Yadav, Sen Pei, Ajitesh Srivastava, Sarah Kefayati, Alexander Rodríguez, Marie-Laure Charpignon, Anil Vullikanti, and B. Aditya Prakash. 2022. epiDAMIK 5.0: The 5th International Workshop on Epidemiology Meets Data Mining and Knowledge Discovery. In *Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*. 4850–4851.
- [4] Aniruddha Adiga, Devdatt Dubhashi, Bryan Lewis, Madhav Marathe, Srinivasan Venkatramanan, and Anil Vullikanti. 2020. Mathematical models for covid-19 pandemic: A comparative analysis. *J. Indian Inst. Sci.* 100, 4 (2020), 793–807.
- [5] Yongsu Ahn and Yu-Ru Lin. 2019. Fairsight: Visual analytics for fairness in decision making. *IEEE Trans. Visual. Comput. Graph.* 26, 1 (2019), 1086–1095.
- [6] Hossein Amiri, Shiyang Ruan, Joon-Seok Kim, Hyunjee Jin, Hamdi Kavak, Andrew Crooks, Dieter Pfoser, Carola Wenk, and Andreas Züfle. 2023. Massive trajectory data based on patterns of life (data and resources paper). In *Proceedings of the 31st ACM SIGSPATIAL International Conference on Advances in Geographic Information Systems*.
- [7] Taylor Anderson and Suzana Dragičević. 2020. NEAT approach for testing and validation of geospatial network agent-based model processes: Case study of influenza spread. *Int. J. Geogr. Info. Sci.* 34, 9 (2020), 1792–1821.
- [8] Taylor Anderson, Amira Roess, Joon-Seok Kim, and Andreas Züfle. 2022. *SpatialEpi'22: Proceedings of the 3rd ACM SIGSPATIAL International Workshop on Spatial Computing for Epidemiology*. Association for Computing Machinery, Seattle, Washington.
- [9] Taylor Anderson, Jia Yu, Amira Roess, Hamdi Kavak, Joon-Seok Kim, and Andreas Züfle. 2023. *SpatialEpi'23: Proceedings of the 4th ACM SIGSPATIAL International Workshop on Spatial Computing for Epidemiology*. Association for Computing Machinery, Hamburg, Germany.
- [10] Taylor Anderson, Jia Yu, and Andreas Züfle. 2021. The 1st ACM SIGSPATIAL International Workshop on Modeling and Understanding the Spread of COVID-19. *SIGSPATIAL Spec.* 12, 3 (2021), 35–40.
- [11] Andreou Andreas, Constandinos X. Mavroumoustakis, George Mastorakis, Shahid Mumtaz, Jordi Mongay Batalla, and Evangelos Pallis. 2020. Modified machine learning Technique for curve fitting on regression models for COVID-19 projections. In *Proceedings of the IEEE 25th International Workshop on Computer Aided Modeling and Design of Communication Links and Networks (CAMAD'20)*. IEEE, 1–6.

- [12] Christopher Antoun, Chan Zhang, Frederick G. Conrad, and Michael F. Schober. 2016. Comparisons of online recruitment strategies for convenience samples: Craigslist, Google AdWords, Facebook, and Amazon Mechanical Turk. *Field Methods* 28, 3 (2016), 231–246.
- [13] Apple. [n.d.]. COVID-19 Mobility Trends Reports. Retrieved from <https://covid19.apple.com/mobility>
- [14] Alexei Baevski, Yuhao Zhou, Abdelrahman Mohamed, and Michael Auli. 2020. wav2vec 2.0: A framework for self-supervised learning of speech representations. *Adv. Neural Info. Process. Syst.* 33 (2020), 12449–12460.
- [15] Eugene Bagdasaryan, Omid Poursaeed, and Vitaly Shmatikov. 2019. Differential privacy has disparate impact on model accuracy. *Adv. Neural Info. Process. Syst.* 32 (2019).
- [16] David Baum. 2008. Reading a phylogenetic tree: The meaning of monophyletic groups. *Nature Edu.* 1, 1 (2008), 190.
- [17] Sumon Biswas and Hriday Rajan. 2021. Fair preprocessing: Towards understanding compositional fairness of data transformers in machine learning pipeline. Retrieved from <https://arXiv:2106.06054>
- [18] Pierre-Yves Boëlle, Cécile Souty, Titouan Launay, Caroline Guerrisi, Clément Turbelin, Sylvie Behillil, Vincent Enouf, Chiara Poletto, Bruno Lina, Sylvie van Der Werf et al. 2020. Excess cases of influenza-like illnesses synchronous with coronavirus disease (COVID-19) epidemic, France, March 2020. *Eurosurveillance* 25, 14 (2020), 2000326.
- [19] GE Box. 1979. All models are wrong, but some are useful. *Robust. Stat.* 202, 1979 (1979), 549.
- [20] Logan C. Brooks, David C. Farrow, Sangwon Hyun, Ryan J. Tibshirani, and Roni Rosenfeld. 2018. Nonmechanistic forecasts of seasonal influenza with iterative one-week-ahead distributions. *PLoS Comput. Biol.* 14, 6 (2018), e1006134.
- [21] Giuseppe C. Calafiore, Carlo Novara, and Corrado Possieri. 2020. A modified SIR model for the COVID-19 contagion in Italy. In *Proceedings of the 59th IEEE Conference on Decision and Control (CDC'20)*. IEEE, 3889–3894.
- [22] Flavio P. Calmon, Dennis Wei, Bhanukiran Vinzamuri, Karthikeyan Natesan Ramamurthy, and Kush R. Varshney. 2017. Optimized pre-processing for discrimination prevention. In *Proceedings of the 31st International Conference on Neural Information Processing Systems*. 3995–4004.
- [23] Yang Cao, Yonghui Xiao, Li Xiong, and Liquan Bai. 2019. PriSTE: From location privacy to spatiotemporal event privacy. In *Proceedings of the IEEE 35th International Conference on Data Engineering (ICDE'19)*. IEEE, 1606–1609.
- [24] Yang Cao, Masatoshi Yoshikawa, Yonghui Xiao, and Li Xiong. 2018. Quantifying differential privacy in continuous data release under temporal correlations. *IEEE Trans. Knowl. Data Eng.* 31, 7 (2018), 1281–1295.
- [25] Alessandro Castelnovo, Riccardo Crupi, Greta Greco, Daniele Regoli, Ilaria Giuseppina Penco, and Andrea Claudio Cosentini. 2022. A clarification of the nuances in the fairness metrics landscape. *Sci. Rep.* 12, 1 (2022), 1–21.
- [26] Estee Y. Cramer, Evan L. Ray, Velma K. Lopez, Johannes Bracher, Andrea Brennen, Alvaro J. Castro Rivadeneira, Aaron Gerding, Tilmann Gneiting, Katie H. House, Yuxin Huang et al. 2022. Evaluation of individual and ensemble probabilistic forecasts of COVID-19 mortality in the United States. *Proc. Natl. Acad. Sci. U.S.A.* 119, 15 (2022), e2113561119.
- [27] Sen Cui, Weishen Pan, Changshui Zhang, and Fei Wang. 2020. xOrder: A model agnostic post-processing framework for achieving ranking fairness while maintaining algorithm utility. Retrieved from <https://arXiv:2006.08267>
- [28] Rachel Cummings, Damien Desfontaines, David Evans, Roxana Geambasu, Yangsibo Huang, Matthew Jagielski, Peter Kairouz, Gautam Kamath, Sewoong Oh, Olga Ohrimenko, and others. 2024. Advancing differential privacy: Where we are now and future directions for real-world deployment. *Harvard Data Science Review* 6, 1 (2024).
- [29] Rachel Cummings, Varun Gupta, Dhamma Kimpara, and Jamie Morgenstern. 2019. On the compatibility of privacy and fairness. In *Proceedings of the 27th Conference on User Modeling, Adaptation and Personalization*. 309–315.
- [30] Demetre Daskalakis, R. Paul McClung, Leandro Mena, Jonathan Mermin, Centers for Disease Control, and Prevention's Monkeypox Response Team. 2022. *Monkeypox: Avoiding the Mistakes of Past Infectious Disease Epidemics*. 1177–1178.
- [31] Shohreh Deldari, Daniel V. Smith, Hao Xue, and Flora D. Salim. 2021. Time series change point detection with self-supervised contrastive predictive coding. In *Proceedings of the Web Conference*. 3124–3135.
- [32] Songgaojun Deng, Shusen Wang, Huzefa Rangwala, Lijing Wang, and Yue Ning. 2020. Cola-GNN: Cross-location attention based graph neural networks for long-term ILI prediction. In *Proceedings of the 29th ACM International Conference on Information and Knowledge Management*. 245–254.
- [33] Descartes Labs. [n.d.]. Data for Mobility Changes in Response to COVID-19. Retrieved from <https://github.com/descarteslabs/DL-COVID-19>
- [34] Catherine Dodds and Ibidun Fakoya. 2020. Covid-19: ensuring equality of access to testing for ethnic minorities. *BMJ (Clinical research ed.)* 369 (2020), m2122.
- [35] E. Dong, H. Du, and L. Gardner. 2020. An interactive web-based dashboard to track COVID-19 in real time. *Lancet Infect. Dis.* 20, 5 (2020), 533–534. [https://doi.org/10.1016/S1473-3099\(20\)30120-1](https://doi.org/10.1016/S1473-3099(20)30120-1)
- [36] Zhanwei Du, Abhishek Pandey, Yuan Bai, Meagan C. Fitzpatrick, Matteo Chinazzi, Ana Pastore y Piontti, Michael Lachmann, Alessandro Vespignani, Benjamin J. Cowling, Alison P. Galvani, and Lauren Ancel Meyers. 2021. Comparative cost-effectiveness of SARS-CoV-2 testing strategies in the USA: A modelling study. *Lancet Public Health* 6, 3 (2021), e184–e191.

- [37] Cynthia Dwork. 2006. Differential privacy. In *Proceedings of the International Colloquium on Automata, Languages, and Programming*. Springer, 1–12.
- [38] Ceyhun Eksin, Keith Paarporn, and Joshua S. Weitz. 2019. Systematic biases in disease forecasting—the role of behavior change. *Epidemics* 27 (2019), 96–105.
- [39] Justin Elarde, Joon-Seok Kim, Hamdi Kavak, Andreas Züfle, and Taylor Anderson. 2021. Change of human mobility during COVID-19: A United States case study. Retrieved from <https://arXiv:2109.09022>
- [40] Claes Enøe, Marios P. Georgiadis, and Wesley O. Johnson. 2000. Estimation of sensitivity and specificity of diagnostic tests and disease prevalence when the true disease state is unknown. *Prevent. Vet. Med.* 45, 1-2 (2000), 61–81.
- [41] Farhad Farokhi. 2021. Optimal pre-processing to achieve fairness and its relationship with total variation barycenter. Retrieved from <https://arXiv:2101.06811>
- [42] A. Stewart Fotheringham and David W. S. Wong. 1991. The modifiable areal unit problem in multivariate statistical analysis. *Environ. Plan. A* 23, 7 (1991), 1025–1044.
- [43] Foursquare. [n.d.]. COVID-19 Foot Traffic Data. Retrieved from <https://aws.amazon.com/marketplace/pp/COVID-19-Foot-Traffic-Data-Free/prodview-cjhkgxpn6vcce>
- [44] Kokouvi M. Gamado, George Streftaris, and Stan Zachary. 2014. Modelling under-reporting in epidemics. *J. Math. Biol.* 69 (2014), 737–765.
- [45] Song Gao, Jinneng Rao, Yuhao Kang, Yunlei Liang, and Jake Kruse. 2020. Mapping county-level mobility pattern changes in the United States in response to COVID-19. *SIGSpatial Spec.* 12, 1 (2020), 16–26.
- [46] Charles E. Gehlke and Katherine Biehl. 1934. Certain effects of grouping upon the size of the correlation coefficient in census tract material. *J. Amer. Statist. Assoc.* 29, 185A (1934), 169–170.
- [47] Jeremy Ginsberg, Matthew H. Mohebbi, Rajan S. Patel, Lynnette Brammer, Mark S. Smolinski, and Larry Brilliant. 2009. Detecting influenza epidemics using search engine query data. *Nature* 457, 7232 (2009), 1012–1014.
- [48] GISAIID. [n.d.]. Phylodynamics of hMpxV. Retrieved from <https://gisaid.org/hmpxv-phylogeny/>
- [49] Google. [n.d.]. COVID-19 Community Mobility Reports. Retrieved from <https://www.google.com/covid19/mobility/>
- [50] Gareth J. Griffith, Tim T. Morris, Matthew J. Tudball, Annie Herbert, Giulia Mancano, Lindsey Pike, Gemma C. Sharp, Jonathan Sterne, Tom M. Palmer, George Davey Smith et al. 2020. Collider bias undermines our understanding of COVID-19 disease risk and severity. *Nature Commun.* 11, 1 (2020), 5749.
- [51] Xiaolan Gu, Ming Li, Yang Cao, and Li Xiong. 2019. Supporting both range queries and frequency estimation with local differential privacy. In *Proceedings of the IEEE Conference on Communications and Network Security (CNS'19)*. IEEE, 124–132.
- [52] Xiaolan Gu, Ming Li, Yueqiang Cheng, Li Xiong, and Yang Cao. 2020. PCKV: Locally differentially private correlated Key-Value data collection with optimized utility. In *Proceedings of the 29th USENIX Security Symposium (USENIX Security'20)*. 967–984.
- [53] Xiaolan Gu, Ming Li, Li Xiong, and Yang Cao. 2020. Providing input-discriminative protection for local differential privacy. In *Proceedings of the IEEE 36th International Conference on Data Engineering (ICDE '20)*. IEEE, 505–516.
- [54] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. 2016. Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 770–778.
- [55] Shaobo He, Yuexi Peng, and Kehui Sun. 2020. SEIR modeling of the COVID-19 and its dynamics. *Nonlin. Dynam.* 101, 3 (2020), 1667–1680.
- [56] Melanie Henwood. 2020. Care home deaths: The untold and largely unrecorded tragedy of COVID-19. *British Policy and Politics at LSE*. Retrieved from <https://blogs.lse.ac.uk/politicsandpolicy/care-home-deaths-covid19/>
- [57] Wei-Ning Hsu, Benjamin Bolte, Yao-Hung Hubert Tsai, Kushal Lakhotia, Ruslan Salakhutdinov, and Abdelrahman Mohamed. 2021. Hubert: Self-supervised speech representation learning by masked prediction of hidden units. *IEEE/ACM Trans. Audio, Speech, Lang. Process.* 29 (2021), 3451–3460.
- [58] Shayan Jawed, Josif Grabocka, and Lars Schmidt-Thieme. 2020. Self-supervised learning for semi-supervised time series classification. In *Proceedings of the Pacific-Asia Conference on Knowledge Discovery and Data Mining*. Springer, 499–511.
- [59] Toshihiro Kamishima, Shotaro Akaho, and Jun Sakuma. 2011. Fairness-aware learning through regularization approach. In *Proceedings of the IEEE 11th International Conference on Data Mining Workshops*. IEEE, 643–650.
- [60] Timo Kaufmann, Paul Weng, Viktor Bengs, and Eyke Hüllermeier. 2023. A survey of reinforcement learning from human feedback. Retrieved from <https://arXiv:2312.14925>
- [61] William Ogilvy Kermack and Anderson G. McKendrick. 1932. Contributions to the mathematical theory of epidemics. II.—The problem of endemicity. *Proc. Roy. Soc. London. Ser. A* 138, 834 (1932), 55–83.
- [62] Katherine M. Keyes, Melissa Tracy, Stephen J. Mooney, Aaron Shev, and Magdalena Cerdá. 2017. Invited commentary: Agent-based models—bias in the face of discovery. *Amer. J. Epidemiol.* 186, 2 (2017), 146–148.
- [63] J. S. Kim, H. Jin, H. Kavak, O. C. Rouly, A. Crooks, D. Pfoser, C. Wenk, and A. Züfle. 2020. Location-based social network data generation based on patterns of life. In *Proceedings of the 21st IEEE International Conference on Mobile Data Management (MDM'20)*. 158–167. <https://doi.org/10.1109/MDM48529.2020.00038>

- [64] Joon-Seok Kim, Hamdi Kavak, Umar Manzoor, Andrew Crooks, Dieter Pfoser, Carola Wenk, and Andreas Züfle. 2019. Simulating urban patterns of life: A geo-social data generation framework. In *Proceedings of the 27th ACM SIGSPATIAL International Conference on Advances in Geographic Information Systems*. 576–579.
- [65] Michael P. Kim, Amirata Ghorbani, and James Zou. 2019. Multiaccuracy: Black-box post-processing for fairness in classification. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*. 247–254.
- [66] Will Kohn, Hossein Amiri, and Andreas Züfle. 2023. EPIPOL: An epidemiological patterns of life simulation (demonstration paper). In *Proceedings of the 4th ACM SIGSPATIAL International Workshop on Spatial Computing for Epidemiology*. 13–16.
- [67] Catherine Kreatsoulas and S. V. Subramanian. 2018. Machine learning in social epidemiology: Learning from experience. *SSM Popul. Health* 4 (2018), 347.
- [68] R. G. Krishnan, S. Cenci, and L. Bourouiba. 2022. Mitigating bias in estimating epidemic severity due to heterogeneity of epidemic onset and data aggregation. *Ann. Epidemiol.* 65 (2022), 1–14.
- [69] T. Kuchler, D. Russel, and J. Stroebel. 2020. *The Geographic Spread of COVID-19 Correlates with Structure of Social Networks as Measured by Facebook*. Technical Report. CESifo Working Paper.
- [70] Tetsuzo Kuragano and Akira Yamaguchi. 2007. Curve shape modification and fairness evaluation. In *Proceedings of the International Design Engineering Technical Conferences and Computers and Information in Engineering Conference*, Vol. 48078. 459–468.
- [71] Yann LeCun, Yoshua Bengio, and Geoffrey Hinton. 2015. Deep learning. *Nature* 521, 7553 (2015), 436.
- [72] Serin Lee, Zelda B. Zabinsky, Judith N. Wasserheit, Stephen M. Kofsky, and Shan Liu. 2021. COVID-19 pandemic response simulation in a large city: Impact of nonpharmaceutical interventions on reopening society. *Med. Decis. Mak.* 41, 4 (2021), 419–429.
- [73] P. Lemey, A. Rambaut, A. J. Drummond, and M. A. Suchard. 2009. Bayesian phylogeography finds its roots. *PLoS Comput. Biol.* 5, 9 (2009), e1000520. <https://doi.org/10.1371/journal.pcbi.1000520>
- [74] Vasilii N. Leonenko and Sergey V. Ivanov. 2018. Prediction of influenza peaks in Russian cities: Comparing the accuracy of two SEIR models. *Math. Biosci. Eng.* 15, 1 (2018), 209–232.
- [75] Ruixuan Liu, Sepanta Zeighami, Haowen Lin, Cyrus Shahabi, Yang Cao, Shun Takagi, Yoko Konishi, Masatoshi Yoshikawa, and Li Xiong. 2023. Supporting pandemic preparedness with privacy enhancing technology. In *Proceedings of the 5th IEEE International Conference on Trust, Privacy and Security in Intelligent Systems and Applications (TPS-ISA'23)*. IEEE Computer Society, 34–43.
- [76] Pranay K. Lohia, Karthikeyan Natesan Ramamurthy, Manish Bhide, Diptikalyan Saha, Kush R. Varshney, and Ruchir Puri. 2019. Bias mitigation post-processing for individual and group fairness. In *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP'19)*. IEEE, 2847–2851.
- [77] Peter Lynn, Alita Nandi, Violetta Parutis, and Lucinda Platt. 2018. Design and implementation of a high-quality probability sample of immigrants and ethnic minorities: Lessons learnt. *Demogr. Res.* 38 (2018), 513–548.
- [78] Charles M. Macal, Nicholson T. Collier, Jonathan Ozik, Eric R. Tatara, and John T. Murphy. 2018. CHISIM: An agent-based simulation model of social interactions in a large urban area. In *Proceedings of the Winter Simulation Conference (WSC'18)*. IEEE, 810–820.
- [79] Madhav Marathe and Anil Kumar S. Vullikanti. 2013. Computational epidemiology. *Commun. ACM* 56, 7 (2013), 88–96.
- [80] Abraham H. Maslow. 1943. A theory of human motivation. *Psychol. Rev.* 50, 4 (1943), 370.
- [81] Edouard Mathieu, Hannah Ritchie, Lucas Rodés-Guirao, Cameron Appel, Charlie Giattino, Joe Hasell, Bobbie Macdonald, Saloni Dattani, Diana Beltekian, Esteban Ortiz-Ospina, and Max Roser. 2020. Coronavirus pandemic (COVID-19). *Our World in Data*. Retrieved from <https://ourworldindata.org/coronavirus>
- [82] Amanda J. Meadows, Ben Oppenheim, Jaclyn Guerrero, Benjamin Ash, Rinette Badker, Cathine K. Lam, Chris Pardee, Christopher Ngoon, Patrick T. Savage, Vikram Sridharan et al. 2022. Infectious disease underreporting is predicted by country-level preparedness, politics, and pathogen severity. *Health Secur.* 20, 4 (2022), 331–338.
- [83] Mohamed Mokbel, Mahmoud Sakr, Li Xiong, Andreas Züfle, Jussara Almeida, Taylor Anderson, Walid Aref, Gennady Andrienko, Natalia Andrienko, Yang Cao, Sanjay Chawla et al. 2022. Mobility data science (Dagstuhl Seminar 22021). *Dagstuhl Reports* 12, 1 (2022), 1–34. <https://doi.org/10.4230/DagRep.12.1.1>
- [84] Mohamed Mokbel, Mahmoud Sakr, Li Xiong, Andreas Züfle, Jussara Almeida, Walid Aref, Gennady Andrienko, Natalia Andrienko, Yang Cao, Sanjay Chawla et al. 2023. Towards mobility data science (vision paper). Retrieved from <https://arXiv:2307.05717>
- [85] David J. Muscatello, Abrar A. Chughtai, Anita Heywood, Lauren M. Gardner, David J. Heslop, and C. Raina MacIntyre. 2017. Translation of real-time infectious disease modeling into routine public health practice. *Emerg. Infect. Dis.* 23, 5 (2017).
- [86] Phi-Yen Nguyen, Whenayon Simeon Ajisegiri, Valentina Costantino, Abrar A. Chughtai, and C. Raina MacIntyre. 2021. Reemergence of human monkeypox and declining population immunity in the context of urbanization, Nigeria, 2017–2020. *Emerg. Infect. Dis.* 27, 4 (2021), 1007.

- [87] John Pesavento, Andy Chen, Rayan Yu, Joon-Seok Kim, Hamdi Kavak, Taylor Anderson, and Andreas Züfle. 2020. Data-driven mobility models for COVID-19 simulation. In *Proceedings of the 3rd ACM SIGSPATIAL International Workshop on Advances in Resilient and Intelligent Cities*. 29–38.
- [88] Rafael Poyiadzi, Kacper Sokol, Raul Santos-Rodriguez, Tijl De Bie, and Peter Flach. 2020. FACE: Feasible and actionable counterfactual explanations. In *Proceedings of the AAAI/ACM Conference on AI, Ethics, and Society*. 344–350.
- [89] David Pujol, Ryan McKenna, Satya Kuppam, Michael Hay, Ashwin Machanavajjhala, and Gerome Miklau. 2020. Fair decision making using privacy-protected data. In *Proceedings of the Conference on Fairness, Accountability, and Transparency*. 189–199.
- [90] Umair Qazi, Muhammad Imran, and Ferda Ofli. 2020. GeoCoV19: A dataset of hundreds of millions of multilingual COVID-19 tweets with location information. *SIGSPATIAL Spec.* 12, 1 (2020), 6–15.
- [91] Nicholas G. Reich, Logan C. Brooks, Spencer J. Fox, Sasikiran Kandula, Craig J. McGowan, Evan Moore, Dave Osthus, Evan L. Ray, Abhinav Tushar, Teresa K. Yamana et al. 2019. A collaborative multiyear, multimodel assessment of seasonal influenza forecasting in the United States. *Proc. Natl. Acad. Sci. U.S.A.* 116, 8 (2019), 3146–3154.
- [92] Alexander Rodríguez, Harshavardhan Kamarthi, Pulak Agarwal, Javen Ho, Mira Patel, Suchet Sapre, and B. Aditya Prakash. 2022. Data-centric epidemic forecasting: A survey. Retrieved from <https://arXiv:2207.09370>
- [93] Bruce Rosen, Ruth Waitzberg, Avi Israeli, Michael Hartal, and Nadav Davidovitch. 2021. Addressing vaccine hesitancy and access barriers to achieve persistent progress in Israel’s COVID-19 vaccination program. *Israel J. Health Policy Res.* 10, 1 (2021), 1–20.
- [94] Jonathan Rothbaum, Jonathan Eggleston, Adam Bee, Mark Klee, and Brian Mendez-Smith. 2021. Addressing nonresponse bias in the American community survey during the pandemic using administrative data. *United States Census Bureau Working Papers*.
- [95] SafeGraph Inc. [n.d.]. SafeGraph Monthly Pattern Data. Retrieved from <https://docs.safegraph.com/docs/monthly-patterns>
- [96] SafeGraph Inc. [n.d.]. Stopping COVID-19 with New Social Distancing Dataset. Retrieved from <https://www.safegraph.com/blog/stopping-covid-19-with-new-social-distancing-dataset>.
- [97] Reva Schwartz, Apostol Vassilev, Kristen Greene, Lori Perine, Andrew Burt, Patrick Hall et al. 2022. Towards a standard for identifying and managing bias in artificial intelligence. *NIST Spec. Publ.* 1270, 10.6028 (2022).
- [98] Weijing Shang, Liangyu Kang, Guiying Cao, Yaping Wang, Peng Gao, Jue Liu, and Min Liu. 2022. Percentage of asymptomatic infections among SARS-CoV-2 omicron variant-positive individuals: A systematic review and meta-analysis. *Vaccines* 10, 7 (2022), 1049.
- [99] Ben Shneiderman. 2020. Human-centered artificial intelligence: Reliable, safe & trustworthy. *Int. J. Hum.-Comput. Interact.* 36, 6 (2020), 495–504.
- [100] Y. Shu and J. McCauley. 2017. GISAI: Global initiative on sharing all influenza data—From vision to reality. *Euro. Surveill.* 22, 13 (2017). <https://doi.org/10.2807/1560-7917.ES.2017.22.13.30494>
- [101] Jake Silberg and James Manyika. 2019. Notes from the AI frontier: Tackling bias in AI (and in humans). *McKinsey Global Inst.* 1, 6 (2019).
- [102] Edward Small, Wei Shao, Zeliang Zhang, Peihan Liu, Jeffrey Chan, Kacper Sokol, and Flora Salim. 2022. How robust is your fair model? Exploring the robustness of diverse fairness strategies. Retrieved from <https://arXiv:2207.04581>
- [103] Kacper Sokol and Peter A. Flach. 2019. Counterfactual explanations of machine learning predictions: Opportunities and challenges for AI safety. In *Proceedings of the AAAI Workshop on Artificial Intelligence Safety (SafeAI@AAAI’19)*.
- [104] Kacper Sokol, Meelis Kull, Jeffrey Chan, and Flora Dilys Salim. 2022. Fairness and ethics under model multiplicity in machine learning. Retrieved from <https://arXiv:2203.07139>
- [105] Kacper Sokol, Raul Santos-Rodriguez, and Peter Flach. 2022. FAT forensics: A Python toolbox for algorithmic fairness, accountability and transparency. *Softw. Impacts* (2022), 100406.
- [106] Adam Spannaus, Theodore Papamarkou, Samantha Erwin, and J. Blair Christian. 2022. Inferring the spread of COVID-19: The role of time-varying reporting rate in epidemiological modelling. *Sci. Rep.* 12, 1 (2022), 10761.
- [107] Yasir Suhail, Junaid Afzal, and Kshitiz. 2021. Incorporating and addressing testing bias within estimates of epidemic dynamics for SARS-CoV-2. *BMC Med. Res. Methodol.* 21 (2021), 1–13.
- [108] Anish Susarla, Austin Liu, Duy Hoang Thai, Minh Tri Le, and Andreas Züfle. 2022. Spatiotemporal disease case prediction using contrastive predictive coding. In *Proceedings of the 3rd ACM SIGSPATIAL Workshop on Spatial Computing for Epidemiology*.
- [109] Christina P. Tadini, Teresa Gisinger, Alexandra Kautzky-Willer, Karolina Kublickiene, Maria Trinidad Herrero, Valeria Raparelli, Louise Pilote, and Colleen M. Norris. 2020. The influence of sex and gender domains on COVID-19 cases and mortality. *Can. Med. Assoc. J.* 192, 36 (2020), E1041–E1045.
- [110] John P. Thornhill, Sapha Barkati, Sharon Walmsley, Juergen Rockstroh, Andrea Antinori, Luke B. Harrison, Romain Palich, Achyuta Nori, Iain Reeves, Maximillian S. Habibi, and others. 2022. Monkeypox virus infection in humans across 16 countries—April–June 2022. *New England Journal of Medicine* 387, 8 (2022), 679–691.

- [111] Hien To, Cyrus Shahabi, and Li Xiong. 2018. Privacy-preserving online task assignment in spatial crowdsourcing with untrusted server. In *Proceedings of the IEEE 34th International Conference on Data Engineering (ICDE '18)*. IEEE, 833–844.
- [112] Alma Tostmann, John Bradley, Teun Bousema, Wing-Kee Yiek, Minke Holwerda, Chantal Bleeker-Rovers, Jaap Ten Oever, Corianne Meijer, Janette Rahamat-Langendoen, Joost Hopman et al. 2020. Strong associations and moderate predictive value of early symptoms for SARS-CoV-2 test positivity among healthcare workers, the Netherlands, March 2020. *Eurosurveillance* 25, 16 (2020), 2000508.
- [113] Jessica Tyrrell, Jie Zheng, Robin Beaumont, Kathryn Hinton, Tom G. Richardson, Andrew R. Wood, George Davey Smith, Timothy M. Frayling, and Kate Tilling. 2021. Genetic predictors of participation in optional components of UK Biobank. *Nature Commun.* 12, 1 (2021), 886.
- [114] Srinivasan Venkatramanan, Bryan Lewis, Jiangzhuo Chen, Dave Higdon, Anil Vullikanti, and Madhav Marathe. 2018. Using data-driven agent-based models for forecasting emerging infectious diseases. *Epidemics* 22 (2018), 43–49.
- [115] Svitlana Volkova, Ellyn Ayton, Katherine Porterfield, and Courtney D. Corley. 2017. Forecasting influenza-like illness dynamics for military populations using neural networks and social media. *PLoS One* 12, 12 (2017), e0188941.
- [116] Han Wang, Hanbin Hong, Li Xiong, Zhan Qin, and Yuan Hong. 2022. L-SRR: Local differential privacy for location-based services with staircase randomized response. In *Proceedings of the ACM Conference on Computer and Communications Security*.
- [117] David W. S. Wong. 2004. The modifiable areal unit problem (MAUP). In *WorldMinds: Geographical Perspectives on 100 Problems*. Springer, 571–575.
- [118] Sean L. Wu, Andrew N. Mertens, Yoshika S. Crider, Anna Nguyen, Nolan N. Pokpongkiat, Stephanie Djajadi, Anmol Seth, Michelle S. Hsiang, John M. Colford, Art Reingold, Benjamin F. Arnold, Alan Hubbard, and Jade Benjamin-Chung. 2020. Substantial underestimation of SARS-CoV-2 infection in the United States. *Nature Commun.* 11, 1 (2020), 1–10.
- [119] Yuexin Wu, Yiming Yang, Hiroshi Nishiura, and Masaya Saitoh. 2018. Deep learning for epidemiological predictions. In *Proceedings of the 41st International ACM SIGIR Conference on Research and Development in Information Retrieval*. 1085–1088.
- [120] Yonghui Xiao and Li Xiong. 2015. Protecting locations with differential privacy under temporal correlations. In *Proceedings of the 22nd ACM SIGSAC Conference on Computer and Communications Security*. 1298–1309.
- [121] Yiqun Xie, Erhu He, Xiaowei Jia, Weiye Chen, Sergii Skakun, Han Bao, Zhe Jiang, Rahul Ghosh, and Praveen Ravithinam. 2022. Fairness by “Where”: A statistically-robust and model-agnostic bi-level learning framework. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 36. 12208–12216.
- [122] Heng Xu and Nan Zhang. 2022. Implications of data anonymization on the statistical evidence of disparity. *Manage. Sci.* 68, 4 (2022), 2600–2618.
- [123] Roy S. Zawadzki, Cynthia L. Gong, Sang K. Cho, Jan E. Schnitzer, Nadine K. Zawadzki, Joel W. Hay, and Emmanuel F. Drabo. 2021. Where do we go from here? A framework for using susceptible-infectious-recovered models for policy making in emerging infectious diseases. *Value Health* 24, 7 (2021), 917–924.
- [124] Hengshuang Zhao, Jiaya Jia, and Vladlen Koltun. 2020. Exploring self-attention for image recognition. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 10076–10085.
- [125] Andreas Züfle, Carola Wenk, Dieter Pfoser, Andrew Crooks, Joon-Seok Kim, Hamdi Kavak, Umar Manzoor, and Hyunjee Jin. 2023. Urban life: A model of people and places. *Comput. Math. Organiz. Theory* 29 (2023), 20–51.

Received 1 April 2023; revised 26 March 2024; accepted 1 April 2024