

Yea or Nay: ‘AI for Social Good’ in the Public Sector

Yu Lu Liu
yliu624@jh.edu
Johns Hopkins University
Baltimore, Maryland, USA

Ziang Xiao
ziang.xiao@jhu.com
Johns Hopkins University
Baltimore, Maryland, USA

Abstract

With the rise of AI projects in the public sector, more and more members of the general public have started to take a stance on these projects. Are these “AI for Social Good” projects worth the public investment? Should AI even be considered as a solution to problems of public interest? These are AI governance questions that the public may look to answer and have their voices heard. In this work, we propose to adopt the policy analysis process framework to analyze public AI projects. We describe how the framework can be applied in this context and highlight challenges that people may face in the process, which we translate into HCI/AI research opportunities.

CCS Concepts

• **Social and professional topics** → **Government technology policy**.

Keywords

AI for Social Good, Responsible AI, AI governance, Civic Engagement

1 Introduction

In September 2024, California governor Gavin Newsom announced that the state would explore the use of generative artificial intelligence (AI) to address homelessness, a challenging issue in California[20]. This announcement received backlash from many social media users. With comments such as “generate some homeless dude,” and “how could AI possibly help? Just give people houses,” these people urged their government to focus its attention on building affordable housing and introducing rent control policies, instead of funding AI developers[11].

From this example, it is clear that projects aiming to solve problems of public interest, what we may call “AI for social good” projects, could be met with public disapproval. Well intentioned AI projects may also lead to bad outcomes. Which public AI projects are worth public investment? Should the government even consider exploring AI solutions in certain contexts? These are governance questions, as the management and allocation of limited resources (e.g., taxpayers’ money) is an integral part of governance. More

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.
Conference’17, Washington, DC, USA

© 2025 Copyright held by the owner/author(s). Publication rights licensed to ACM.
ACM ISBN 978-x-xxxx-xxxx-x/YYYY/MM
<https://doi.org/10.1145/nnnnnnnn.nnnnnnnn>

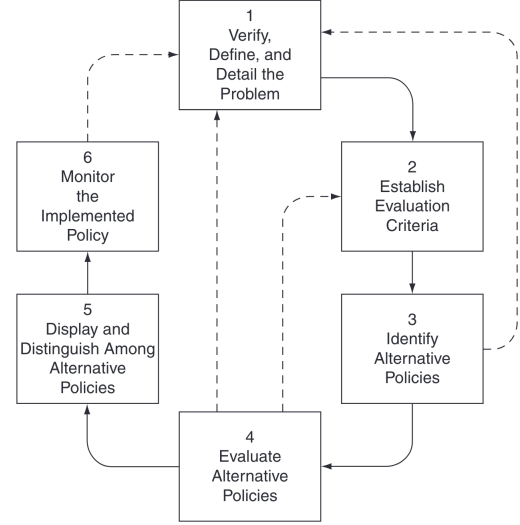


Figure 1: Policy Analysis Process by Patton et al. [23], based on the Classical Rational Problem-Solving Process.

specifically, these are AI governance questions that must not be overlooked. To help answering these questions, we propose to apply the policy analysis process framework proposed by Patton et al. [23] which guides the process of identifying and evaluating policies or programs that are intended to lessen or resolve social, economic, or physical problems. The framework decomposes policy analysis into 6 steps as illustrated in Figure 1: i) verify, define, and detail the problem, ii) establish evaluation criteria, iii) identify alternative solutions, iv) evaluate solutions, v) display and distinguish among solutions, and vi) monitor the implemented solution.

By reasoning through these steps, we surface challenges that people, especially non-expert citizens, may face in the process of analyzing public AI projects. We then translate these challenges into opportunities for HCI and AI researchers interested in supporting this process and in supporting citizen involvement.

2 Related Work

2.1 Policy Analysis and Citizen Involvement

Policy Analysis and Citizen Participation Policy analysis is defined as “the process through which we identify and evaluate policies or programs that are intended to lessen or resolve social, economic, or physical problems.” [23] Policy analysts are often third-party contractors, external to the agency trying to make policy decisions, because it is often believed that they are more objective, therefore leading to better analysis. As a result, policy analysis usually only involve decision-makers and professional analysts that they hire.

This approach to policy analysis has been criticized for not only failing to solve social problems but even contributing to them [19]. Citizen participation to policy analysis is believed to be the better way forward, as it allows the analysis to i) take into account citizen knowledge and ideas on public issues, ii) build public support for final decisions, and iii) build a relationship of collaboration and trust between decision-makers and the public [7]. Unlike for conventional policy analysis, where success is measured by the extent to which objectives are achieved, this approach to policy analysis emphasizes whether balance is found among competing interests and whether consensus is reached on appropriate actions forward.

Citizen participation can come in different forms, which Cogan et al. [6] categorize as follows:

- Publicity: promoting a certain policy or program to persuade and gather public support.
- Public education: present complete and balanced information so that citizens can draw their own conclusions about the policy/program.
- Public input: solicit ideas and opinions from citizens, ideally coupled with a feedback mechanism that informs the public about how their input has shaped the final decision
- Public interaction: citizens, policy analysts, and decision makers interact with each other, allowing for an exchange of ideas that contributes to a consensus.
- Public partnership: citizens have a formalized role in shaping the ultimate decisions, becoming in some way decision-makers themselves.

Citizen participation can also be performative, not actually involving stakeholders in a genuine and meaningful way. Arnstein's "ladder of citizen participation"[1] is a framework that helps to assess how much agency and control is given to citizens:

- (1) Manipulation: people are lied to about e.g., project goals in order to secure public approval;
- (2) Therapy: experts focus on "fixing" people's beliefs;
- (3) Informing: an unidirectional information flow to the public;
- (4) Consultation: a bidirectional information flow, but people's views might not be taken into account;
- (5) Placation: decision-makers create advising committees of citizens with minimal authority;
- (6) Partnership: negotiation between citizens and decision-makers provides citizens with more power;
- (7) Delegated Power: citizens has dominant power;
- (8) Citizen Control: citizens are in complete control of the project.

We will draw from this body of prior work when imagining different ways citizens could be involved in the decision about whether AI-based solution should be implemented to address problems of public interest.

2.2 AI in the Public Sector

AI use in the public sector, often grouped under the terminology of "algorithmic decision-making", spans many domains, such as child welfare, public housing, public health, and law enforcement. Prior work examining public AI have focused on issues around trust. For example, Brown et al. [5] study prediction tools deployed by child welfare agencies to predict child safety risk, suggesting cases for authorities to investigate. They organized workshops to

learn about the concerns of communities affected by the use and development of these tools and outline future directions aiming to raise people's comfort. For example, they recommend future development effort to model "success" factors (e.g., parents having stable income) instead of only focusing on modelling "failure" factors (e.g., parents having criminal records). Drobotowicz et al. [10] seek to answer the same research question, conducting interviews about the use of AI in public services that i) make decisions about access to housing, ii) make mental health prediction, iii) assess education impact on children, or iv) detect financial fraud in social insurance organization. They surface requirements citizens have for trustworthy AI services in the public sector, such as transparency about the AI process in order for citizens to make informed decisions, and justification as to why AI is used in public service (i.e., "What is the reason for using AI in the public service?")

2.3 Participatory AI

Participatory AI generally refers to when participatory design methods are applied to AI development cycle. For example, the training data for an AI model could be labeled by a specific stakeholder group with the goal of creating a model that is aligned to this group's views or preferences.

Like for participatory design in other field, participatory AI also run the risk of being performative. Corbett et al. [8] employed Arnstein's "ladder of citizen participation" framework (Section 2.1) to assess how much agency and control is given to participants in participatory AI projects. Using this framework, the authors analyzed 21 participatory AI papers and found that most work only informs or consults (rung iii and iv on the metaphorical ladder). As an example of an AI project where participants are in complete control (highest rung on the ladder), Nekoto et al. [18] conducted participatory research, where individuals involved in the research don't necessarily have formal research training, to develop machine translation systems for African languages. The project starts with the premise of developing the systems, although this goal stems from participants' needs and not imposed through some external authority. Participatory AI projects tend to focus on *how* to use or develop AI. The decision about *whether or not* to use or develop AI, which is the topic of interest for this work, is generally out of the scope of these projects.

3 Policy Analysis Process for Public AI

We propose to reflect about public AI projects by following the policy analysis process framework proposed by Patton et al. [23] (illustrated in Figure 1). For each step in the framework, we first describe it, then discuss how it applies to analyze public AI projects, and finally outline challenges that people may face at this step, especially focusing on citizen involvement.

Step 1. Verify, define, and detail the problem.

The first step in the policy analysis framework is to characterize the problem that we are trying to solve, which includes i) verifying whether the problem actually exists, ii) determining the extent and magnitude of the problem and iii) determining the stakes and the stakeholders involved. We believe that this must also be the first step when analyzing proposals of using AI to solve a problem. Under the

view that AI should not be considered the unquestioned solution to every problem, it is important to reflect about the compatibility between AI and the problem at hand: what about the problem that makes AI a promising solution? There should be a clear justification as to why AI is being considered as a solution in the first place.

To answer this question, we also need clarity about the "AI" that is being considered as a potential solution: does the project involve prediction based on historical data, or does it only deliver insights about the data? Is it generative AI deployed on government employees' computers, or AI-powered robots automating physical tasks? Different methods have different strengths and limitations, yet "AI" is often an overloaded concept that obfuscates those differences. It is thus important to have transparency about the exact methodology being proposed.

Challenges and Opportunities. Citizens funding and being impacted by the potential use of public AI systems are naturally stakeholders who should be involved in characterizing the problem that the systems aim to solve. Ideally, the involvement would take the form of a "public interaction" where citizens and decision-makers can arrive at a consensus via exchange of views and ideas. Before that, however, there is an initial hurdle that is the lack of knowledge and engagement about the problem at hand.

This challenge is common to the analysis of any public project, regardless of whether AI is involved: e.g., the public may not know nor care enough about the issue of homelessness in their community. Additionally, specifically for when AI is involved, the public may not know enough about the "AI" that is being considered as a solution. Improving civic education (including AI literacy) would help overcome these challenges.

In HCI, the field of digital civics has been studying how to better support civic education and engagement through digital technologies [16, 22, 24, 29]. For example, Peacock et al. [24] explored using digitally-supported walks and online discussion platforms in order to increase civic engagement of children in urban design. Similar efforts exist for AI literacy [4, 15]: Lee et al. [13] explored using online workshop sessions so that middle school students can become informed citizens and critical users of AI.

Step 2. Establish evaluation criteria.

The second step is to establish criteria that allow us to compare between proposed solutions or determine when a proposed solution is acceptable. These criteria are then to be operationalized as measures, which are used to evaluate and select amongst proposed solutions. Policy analysts commonly measure criteria such as cost, effectiveness, administrative ease, legality, and political acceptability. It is also important to identify which criteria are the most relevant to various stakeholder groups of the project, so to determine which criteria are central in the analysis. This step is done early in the analysis process so to avoid rationalizing preferred solutions (i.e., picking the solution and then conducting evaluation in a way that justifies it).

The established evaluation criteria should cover intended benefits, as well as costs and risks. When it comes to intended benefits, they are often the most advertised aspect of a project, but they can be vaguely described, making it difficult to later verify whether the project achieved its goal. For example, instead of simply "this

project aims to reduce the homeless population in California", we should obtain more clarity about the exact goal: by how much do we want to reduce it by? How exactly would the homeless population be benefiting? As for costs, for an AI project, we could aim to measure criteria such as computational resources required to develop the system, as well as the environmental impacts of running the system. Finally, for risks, a myriad of potential negative impacts could be relevant: privacy risks (e.g., to people whose data is used for training or inference), social biases being perpetuated (e.g., by system performing prediction based on historical data), and more.

Challenges and Opportunities. This step involves many choices that encode people's values. For example, we may aim to measure biases in system decisions because we value fairness. As a result, it is important to make sure that the system evaluation reflects the values and the concerns of the public (i.e., at least requiring "public input"). Here, we could adopt and improve participatory design practices such as value-sensitive design for AI systems [25, 30]. One core component of value-sensitive design involves conducting conceptual investigations: eliciting values from stakeholders and conceptualizing them to later conduct empirical evaluation.

Step 3. Identify alternative solutions.

Having characterized the problem and identified relevant evaluation criteria, policy analysts are better prepared to find or create alternative solutions to consider in the analysis. Note that inaction, i.e., maintaining the status-quo, is sometimes worth being considered as a promising "solution".

In the motivating example of using AI to address homelessness, many social media users urged the government to build shelters and provide affordable housing through stricter rent control — solutions that have been long advocated for and that are not centered around AI. If such alternative solutions exist, or even better, are shown to work, then they should not be overlooked by governments in order to blindly pursue AI solutions. In mapping out the space of solutions, whether they involve AI or not, we can also better compare them against each other and make a more sensible judgment.

Opportunities. This is not to say that there is a dichotomy between AI and these non-AI alternative solutions. Quite the opposite: they can work in concert. For example, Umbrello and van de Poel [26] developed a machine learning system to detect tenants vulnerable to landlord harassment in New York City, so to prioritize the city's outreach efforts to inform tenants of their rights and assist them. The main policy is the rent-stabilization policies that the city adopted, with the system playing a supportive role. Exploring alternative solutions can thus be a source of inspiration and opportunities for HCI researchers and AI practitioners: instead of "solving" a problem with AI, can we better support existing solutions?

Step 4. Evaluate solutions.

The most vital part of the framework is to evaluate solutions: to what extent does each candidate solution satisfy the evaluation criteria defined in Step 2? This requires operationalizing these criteria. It is the nature of the problem and the specified evaluation criteria that will inform the choice of evaluation methods: policy

analysts are encouraged to avoid the “toolbox approach” where they apply their favourite (or most familiar) method to tackle any criteria. When it comes to evaluating AI systems, negligence in evaluation risks resulting in the deployment of systems that are ineffective and even harmful.

We see two possible scenarios. First, some evaluation may have already been conducted. For example, when AI developers present existing products to be purchased for public use, they might show existing evaluation results and information about their evaluation process. We must then determine whether these results support the effectiveness and safety of the AI system. These existing results may not generalize to cover the specific context of the problem at hand, especially if these results come from “general-purpose” benchmarks, commonplace in the AI field. Second, assuming that there is no existing evaluation results that are informative and useful for our purposes, we must conduct contextualized evaluation.

Challenges and Opportunities. The above two scenarios converge to one thing: the need to improve AI evaluation practices to deliver useful information about the performance of AI system in specific contexts. AI evaluation is an active area of research. One particular challenge is the aforementioned difficulty in mapping evaluation results from established methods (e.g., “general-purpose” benchmarks) to systems’ real world behaviors and impacts [14, 27]. Many initiatives aim to bridge the field of HCI and AI-related fields, for instance, the field of natural language processing [3, 28]. These initiatives encourage the development of evaluation methods that better reflect systems’ real world performance, by taking inspiration from evaluation practices in HCI. These research communities could thus provide great opportunities for HCI researchers to collaborate with AI practitioners.

Public auditing interfaces could also better involve the public in system evaluation by allowing people to examine and critique outputs of AI systems deployed in their communities [9, 12]. For example, Lam et al. [12] designed a public auditing framework where end-users label a small amount of system outputs, which are used to create personalized metrics to conduct a complete audit. This line of work facilitates the process of auditing, possibly allowing for citizens to conduct audits independently.

Diversity in expertise. Note that some criteria of interest might be unfamiliar to AI and HCI experts, or out of the scope of computer science altogether. For example, if political acceptability is a key criteria specified in Step 2, then this current step may require experts from political science. Needless to say, the evaluation methods for alternative solutions most likely require other expertise.

Step 5-6. Select and monitor the implemented solution.

Steps following the evaluation involves presenting the evaluation results so decision-makers can select the solution to implement, and to finally monitor the implemented solution to make sure that it achieves the intended goal. What are plans to ensure that the implemented solution has and keeps having the intended positive impacts that it aims to have, to decide when and whether the solution should be paused, modified, or terminated? Negligence in

this last stage makes it so that the public cannot hold decision-makers and AI developers accountable for the broader impacts of the project.

Current efforts to track the impacts of deployed systems include BLIP [21] and the AI Incident Database [17]. BLIP (which is not specific to AI) uses online articles reporting negative impacts of technology and categorizes them in an interactive web interface. AI Incident Database groups publicly available incident reports (documents from the popular, trade, and academic press), acting as a “collective memory” of AI development and deployment failures. Future work in this direction could explore how to track public AI (e.g., collecting information from government documents) and whether/why implemented AI solution succeeded or failed.

4 Conclusion

Discussion surrounding AI governance often revolves around laws and other regulatory measures. In this work, we bring attention to an aspect of governance that can be often neglected — public projects involving AI. We propose to adopt existing framework in policy analysis to reflect about these projects, and suggest some HCI research directions that could support this process.

We have focused on highlight technical solutions, but we acknowledge that significant institutional changes may be needed in order to have genuine citizen participation in decisions surrounding public AI, for example, when local councils are uncooperative and unwilling to involve citizens in decisions (regardless of whether AI is involved). In this scenario, citizens could choose to engage in “illegitimate” civic participation (i.e., activism outside of formal political and institutional channels). Information communication technologies play a significant role in shaping this form of civic participation [2] — a promising direction for HCI researchers. Our work is further limited by our assumption that public AI projects stem from good intentions in the first place (i.e. “AI4SG”). We must contend with the fact that governments may have ulterior motives and use the label of “AI4SG” as a shield against criticisms. It may be worth examining what are the harms of labelling a project so.

References

- [1] Sherry R Arnstein. 1969. A ladder of citizen participation. *Journal of the American Institute of planners* 35, 4 (1969), 216–224.
- [2] Mariam Asad and Christopher A. Le Dantec. 2015. Illegitimate Civic Participation: Supporting Community Activists on the Ground. In *Proceedings of the 18th ACM Conference on Computer Supported Cooperative Work & Social Computing* (Vancouver, BC, Canada) (CSCW '15). Association for Computing Machinery, New York, NY, USA, 1694–1703. doi:10.1145/2675133.2675156
- [3] Su Lin Blodgett, Amanda Cercas Curry, Sunipa Dev, Michael Madaio, Ani Nenkova, Diyi Yang, and Ziang Xiao (Eds.). 2024. *Proceedings of the Third Workshop on Bridging Human-Computer Interaction and Natural Language Processing*. Association for Computational Linguistics, Mexico City, Mexico. <https://aclanthology.org/2024.hcinlp-1.0/>
- [4] Jason Borenstein and Ayanna Howard. 2021. Emerging challenges in AI and the need for AI ethics education. *AI and Ethics* 1, 1 (Feb. 2021), 61–65.
- [5] Anna Brown, Alexandra Chouldechova, Emily Putnam-Hornstein, Andrew Tobin, and Rhema Vaithianathan. 2019. Toward Algorithmic Accountability in Public Services: A Qualitative Study of Affected Community Perspectives on Algorithmic Decision-making in Child Welfare Services. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (Glasgow, Scotland UK) (CHI '19). Association for Computing Machinery, New York, NY, USA, 1–12. doi:10.1145/3290605.3300271
- [6] Arnold Cogan, Sumner Sharpe, and Joe Hertzberg. 1986. Citizen participation. *The practice of state and regional planning* 290 (1986).
- [7] C Cogan and G Sharpe. 1986. Planning analysis: The theory of citizen involvement. *The Theory of Citizen Participation, University of Oregon* (1986), 284–298.

- [8] Eric Corbett, Emily Denton, and Sheena Erete. 2023. Power and Public Participation in AI. In *Proceedings of the 3rd ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization* (Boston, MA, USA) (EAAMO '23). Association for Computing Machinery, New York, NY, USA, Article 8, 13 pages. doi:10.1145/3617694.3623228
- [9] Alicia DeVos, Aditi Dhabalia, Hong Shen, Kenneth Holstein, and Motahhare Eslami. 2022. Toward User-Driven Algorithm Auditing: Investigating users' strategies for uncovering harmful algorithmic behavior. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (New Orleans, LA, USA) (CHI '22). Association for Computing Machinery, New York, NY, USA, Article 626, 19 pages. doi:10.1145/3491102.3517441
- [10] Karolina Drobotowicz, Marjo Kauppinen, and Sari Kujala. 2021. Trustworthy AI Services in the Public Sector: What Are Citizens Saying About It?. In *Requirements Engineering: Foundation for Software Quality*, Fabiano Dalpiaz and Paola Spoletini (Eds.). Springer International Publishing, Cham, 99–115.
- [11] Liam Geraghty. 2024. California's governor wants to use AI to solve homelessness – but it's backfired, badly. *Big Issue* (2024). <https://www.bigissue.com/news/housing/california-governor-gavin-newsom-ai-homelessness/>
- [12] Michelle S. Lam, Mitchell L. Gordon, Danaë Metaxa, Jeffrey T. Hancock, James A. Landay, and Michael S. Bernstein. 2022. End-User Audits: A System Empowering Communities to Lead Large-Scale Investigations of Harmful Algorithmic Behavior. *Proc. ACM Hum.-Comput. Interact.* 6, CSCW2, Article 512 (Nov. 2022), 34 pages. doi:10.1145/3555625
- [13] Irene Lee, Safinah Ali, Helen Zhang, Daniella DiPaola, and Cynthia Breazeal. 2021. Developing Middle School Students' AI Literacy. In *Proceedings of the 52nd ACM Technical Symposium on Computer Science Education* (Virtual Event, USA) (SIGCSE '21). Association for Computing Machinery, New York, NY, USA, 191–197. doi:10.1145/3408877.3432513
- [14] Q. Vera Liao and Ziang Xiao. 2025. Rethinking Model Evaluation as Narrowing the Socio-Technical Gap. arXiv:2306.03100 [cs.HC] <https://arxiv.org/abs/2306.03100>
- [15] Duri Long and Brian Magerko. 2020. What is AI Literacy? Competencies and Design Considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (CHI '20). Association for Computing Machinery, New York, NY, USA, 1–16. doi:10.1145/3313831.3376727
- [16] Anna Paola Marconi, Gianluca Schiavo, Massimo Zancanaro, Giuseppe Valetto, and Marco Pistore. 2018. Exploring the world through small green steps: improving sustainable school transportation with a game-based learning interface. In *Proceedings of the 2018 International Conference on Advanced Visual Interfaces* (Castiglione della Pescaia, Grosseto, Italy) (AVI '18). Association for Computing Machinery, New York, NY, USA, Article 24, 9 pages. doi:10.1145/3206505.3206521
- [17] Sean McGregor. 2020. Preventing Repeated Real World AI Failures by Cataloging Incidents: The AI Incident Database. arXiv:2011.08512 [cs.CY] <https://arxiv.org/abs/2011.08512>
- [18] Wilhelmina Nekoto, Vukosi Marivate, Tshinondiwa Matsila, Timi Fasubaa, Taiwo Fagbohunbe, Solomon Oluwole Akinola, Shamsuddeen Muhammad, Salomon Kabongo Kabenamualu, Salomey Osei, Freshia Sackey, Rubungo Andre Niyongabo, Ricky Macharm, Perez Ogayo, Orevaoghene Ahia, Musie Meressa Berhe, Mofetoluwa Adeyemi, Masabata Mokgesi-Seling, Lawrence Okegbemi, Laura Martinus, Kolawole Tajudeen, Kevin Degila, Kelechi Ogueji, Kathleen Siminyu, Julia Kreutzer, Jason Webster, Jamiil Toure Ali, Jade Abbott, Irore Orife, Ignatius Ezeani, Idris Abdulkadir Dangana, Herman Kamper, Hady Elsahar, Goodness Duru, Ghollah Kioko, Murhabazi Espoir, Elan van Biljon, Daniel Whitenack, Christopher Onyefuluchi, Chris Chinenye Emezue, Bonaventure F. P. Dossou, Blessing Sibanda, Blessing Bassey, Ayodele Olabiye, Arshath Ramkilowan, Alp Öktem, Adewale Akinfaderin, and Abdallah Bashir. 2020. Participatory Research for Low-resourced Machine Translation: A Case Study in African Languages. In *Findings of the Association for Computational Linguistics: EMNLP 2020*, Trevor Cohn, Yulan He, and Yang Liu (Eds.). Association for Computational Linguistics, Online, 2144–2160. doi:10.18653/v1/2020.findings-emnlp.195
- [19] Dorothy Nelkin. 1981. *Nuclear Power and Its Critics; the Cayuga Lake Controversy*. New York: George Braziller.
- [20] Gavin Newsom. 2024. Governor Newsom seeks to harness the power of GenAI to address homelessness, other challenges. *Governor Gavin Newsom Newsroom* (2024). <https://www.gov.ca.gov/2024/09/05/governor-newsom-seeks-to-harness-the-power-of-genai-to-address-homelessness-other-challenges/>
- [21] Rock Yuren Pang, Sebastin Santy, René Just, and Katharina Reinecke. 2024. BLIP: Facilitating the Exploration of Undesirable Consequences of Digital Technologies. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (CHI '24). Association for Computing Machinery, New York, NY, USA, Article 290, 18 pages. doi:10.1145/3613904.3642054
- [22] Irina Paraschivoiu, Josef Buchner, Robert Praxmarer, and Thomas Layer-Wagner. 2021. Escape the Fake: Development and Evaluation of an Augmented Reality Escape Room Game for Fighting Fake News. In *Extended Abstracts of the 2021 Annual Symposium on Computer-Human Interaction in Play* (Virtual Event, Austria) (CHI PLAY '21). Association for Computing Machinery, New York, NY, USA, 320–325. doi:10.1145/3450337.3483454
- [23] Carl Patton, David Sawicki, and Jennifer Clark. 2015. *Basic methods of policy analysis and planning*. Routledge.
- [24] Sean Peacock, Robert Anderson, and Clara Crivellaro. 2018. Streets for People: Engaging Children in Placemaking Through a Socio-technical Process. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems* (Montreal QC, Canada) (CHI '18). Association for Computing Machinery, New York, NY, USA, 1–14. doi:10.1145/3173574.3173901
- [25] Steven Umbrello and Ibo van de Poel. 2021. Mapping value sensitive design onto AI for social good principles. *AI and Ethics* 1, 3 (Aug. 2021), 283–296.
- [26] Steven Umbrello and Ibo van de Poel. 2021. Mapping value sensitive design onto AI for social good principles. *AI and Ethics* 1, 3 (Aug. 2021), 283–296.
- [27] Kiri L. Wagstaff. 2012. Machine learning that matters. In *Proceedings of the 29th International Conference on Machine Learning* (Edinburgh, Scotland) (ICML '12). Omnipress, Madison, WI, USA, 1851–1856.
- [28] Ziang Xiao, Wesley Hanwen Deng, Michelle S. Lam, Motahhare Eslami, Juho Kim, Mina Lee, and Q. Vera Liao. 2024. Human-Centered Evaluation and Auditing of Language Models. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (CHI EA '24). Association for Computing Machinery, New York, NY, USA, Article 476, 6 pages. doi:10.1145/3613905.3636302
- [29] Wei Zheng, Yan Zhou, and Yi Qin. 2019. An Empirical Study of Incorporation of Augmented Reality into Civic Education. In *Proceedings of the 2019 International Conference on Modern Educational Technology* (Nanjing, China) (ICMET 2019). Association for Computing Machinery, New York, NY, USA, 30–34. doi:10.1145/3341042.3341054
- [30] Haiyi Zhu, Bowen Yu, Aaron Halfaker, and Loren Terveen. 2018. Value-Sensitive Algorithm Design: Method, Case Study, and Lessons. *Proc. ACM Hum.-Comput. Interact.* 2, CSCW, Article 194 (Nov. 2018), 23 pages. doi:10.1145/3274463