

---

# Empowering Users Together: Connecting Algorithmic Collective Action and Explainable AI

---

**Ayana Hussain**

Department of Computer Science  
Simon Fraser University  
Burnaby, BC Canada  
ayana\_hussain@sfu.ca

**Cole Michael Thacker**

Department of Computer Science  
Simon Fraser University  
Burnaby, BC Canada  
cta127@sfu.ca

**Patrick Zhao**

Department of Computer Science  
Simon Fraser University  
Burnaby, BC Canada  
patrick\_zhao@sfu.ca

**Aditya Karan**

Department of Computer Science  
University of Illinois  
Urbana-Champaign, US  
karan2@illinois.edu

**Nicholas Vincent**

Department of Computer Science  
Simon Fraser University  
Burnaby, BC Canada  
nvincent@sfu.ca

## Abstract

1       Algorithmic Collective Action (ACA) and Explainable AI (XAI) both aim to  
2       empower users, yet approach this goal differently: XAI tends to explain “why”  
3       algorithms make a prediction, while ACA tends to focus on “how” users can  
4       collectively influence outcomes. Despite their similar goals and challenges, there  
5       is a notable lack of research that addresses or connects both fields. Another related  
6       field, algorithmic recourse research, highlights this intersection but remains mostly  
7       limited to individual interventions. In this paper, we analyze conceptual overlaps  
8       between XAI and ACA, identify practical opportunities for integration, and provide  
9       examples showing how XAI’s explanatory methods can enhance ACA strategies  
10      while ACA work can inform more actionable XAI design. Our findings support that  
11      combining explanation and action can more effectively advance user understanding,  
12      agency, and algorithmic transparency.

## 13   1 Introduction

14   Digital platforms use machine learning (ML) algorithms to make decisions that affect millions of  
15   users, from content recommendation to loan approvals [30, 18]. As these systems spread, two research  
16   communities have emerged to address the growing need for algorithmic accountability: Explainable  
17   Artificial Intelligence (XAI) and Algorithmic Collective Action (ACA). XAI has grown into a well-  
18   established field with extensive research, open-source tools, and deep industry involvement [9]. XAI  
19   research focuses on making AI systems more transparent and interpretable to users, developers,  
20   and regulators [6, 21], while ACA is a nascent field that examines how groups of users influence  
21   algorithmic systems to achieve a shared goal [26, 12].

22 However, given the emerging nature of ACA, there is limited research connecting it to XAI, and  
23 thus far, the two fields have developed largely in isolation despite addressing complementary aspects  
24 of human-algorithm interaction. This separation creates a missed opportunity, as both fields face  
25 similar challenges. For example, prior ACA work interviewed fans who led large-scale collective  
26 actions [30], with many emphasizing that collective action will fail without a clear understanding of  
27 how a platform’s algorithms function. XAI can offer simplified explanations of how these systems  
28 work, which could help users build better intuitions about effective collective action. This may  
29 be particularly helpful as fans often understand that likes, comments, and shares may influence  
30 recommendation outcomes, but remain uncertain about the downstream effects of inputs on ranking  
31 or visibility [30].

32 XAI research often seeks to provide explanations that are actionable [25]; however, explanations  
33 are often static [31] and difficult for non-technical users to interpret or act on [3]. Previous work  
34 has called for more user-centered, directive, and interactive forms of explanations that support both  
35 understanding and actionable recourse [31, 3]. Thus, studying ACA strategies and designing XAI  
36 interfaces for similar collective and goal-driven scenarios could reveal concrete design requirements  
37 for more actionable explanations.

38 The goal of this short paper is to argue that the XAI and ACA communities often share fundamental  
39 goals of empowering users to engage strategically with public-facing algorithms, and hence should  
40 work together to more fully support user understanding, empowerment, and participation in algorithmic  
41 systems. XAI empowers ACA by providing proven methods for understanding ML systems,  
42 which enables more effective engagement with digital platforms that use such systems. In turn, ACA  
43 reveals both the utility and limitations of XAI methods when communities coordinate to understand  
44 and influence these platforms.

45 The structure of the paper is as follows: Section 2 reviews XAI, data valuation, and ACA; Sections  
46 3-6 present our position, common ground, counterarguments, and recommendations, comparing data  
47 valuation and XAI in relation with ACA; Section 7 lists future work, and Section 8 concludes with  
48 key takeaways and a vision for empowering users through collaborative work across these areas.

## 49 **2 Background**

### 50 **2.1 Explainable AI and Algorithmic Recourse**

51 Explainable AI addresses the black-box nature of ML systems. XAI aims to provide human-  
52 understandable justification for algorithmic decisions, with motivations such as improving trust-  
53 worthiness, enhancing transparency, and promoting fairness [6].

54 Recent research has emphasized the need for user-centered approaches to XAI, moving beyond  
55 algorithmic-centered explanations toward those that meet actual human information needs. For  
56 example, Ehsan et al. highlighted gaps between XAI algorithmic output and the properties of  
57 explanations that individuals actually seek [7], while other work identified disconnects between XAI  
58 outputs and human reasoning processes [31].

59 The field has also struggled with actionability [25]. While traditional XAI focuses on explaining  
60 the “why” behind decisions, researchers increasingly recognize the need for “how” explanations that  
61 guide users with actionable steps towards their goals [18, 25].

62 Algorithmic recourse aims to address this gap by providing explanations and recommendations for  
63 individuals to change unfavorable algorithmic decisions [18]. However, current algorithmic recourse  
64 work typically focuses on individual-level interventions with two main approaches to providing  
65 recourse. Firstly, contrastive recourse/explanations, which suggest feature changes to flip undesired  
66 decisions, and secondly, causal recourse or consequential recommendations, where causal models  
67 generate feasible, real-world interventions accounting for downstream effects [8, 17].

68 Recourse can also involve proxy action, such as family members or a legal representative acting on  
69 someone’s behalf, and collective action, which remains underexplored [17] (though some work has  
70 begun to connect ACA with recourse [5]). Similarly, algorithmic recourse faces several challenges,  
71 centered around addressing: limitations of realism (alignment between recommended action and  
72 real-world possibilities), changes in the environment over time, competition among individuals  
73 pursuing similar interventions, and the recipient’s ability or willingness to act [8].

74 Although most work centers on individual level recourse, a smaller line of work has explored recourse  
 75 from a group or subgroup level. Gupta et al., for example, formalize recourse at the group level  
 76 by measuring the distance from an individual to the classifier decision boundary, and introduce a  
 77 regularized objective that minimizes differences in recourse across groups [11]. Similarly, Rawal and  
 78 Lakkaraju, develop a framework that learns global population level summaries of actionable recourse  
 79 and can generate subgroup specific recourse rules (based on either user specified or automatically  
 80 discovered subgroups) [23]. This allows decision makers to examine how recourse differs across  
 81 subgroups and to identify potential biases or discriminatory outcomes. Ultimately, these works  
 82 highlight the value of examining recourse beyond strictly individual levels, and the opportunity for  
 83 collectives to use XAI and recourse to understand and address group level disparities.

84 As we will describe below, there is great potential for direct action by collectives to “offload” or  
 85 delegate the power for achieving recourse directly to users. Pairing explanations with easy-to-access  
 86 ACA could enable many of the outcomes recourse strategies aim to achieve [5].

## 87 2.2 Data Valuation

88 Data valuation investigates how individual or groups of data points contribute to an ML model’s  
 89 prediction. One such method, influence functions, estimates how a model’s output would change if  
 90 a training point were removed or perturbed without requiring retraining [19]. Other data valuation  
 91 approaches include the use of Shapley values, which fairly distribute credit among training examples  
 92 based on their contribution to the overall model performance [10]. This valuation could enable the  
 93 distribution of rewards generated by large AI models trained on public data, or help data holders  
 94 understand their contribution to said models. Data valuation could also help guide effective ACA,  
 95 as it provides the quantitative foundation on which groups can build claims for compensation or  
 96 redistribution, or determine who holds leverage in a negotiation.

97 Although prior works have explored data valuation methods in adversarial contexts such as data  
 98 poisoning [15], ACA may adopt similar strategies for a spectrum of intents to influence model  
 99 behavior, ranging from constructive efforts at value alignment to intentionally harmful or disruptive  
 100 interventions. Studying these methods in the context of ACA can help clarify how collectives, both  
 101 benevolent and harmful, might leverage data influence and inform responses to both legitimate and  
 102 malicious uses.

103 Fundamentally, data valuation techniques produce quantitative values that can support explanation.  
 104 Most techniques make statements about the causal impact of data, and many XAI methods even use  
 105 concepts from game theory that also appear in technical data valuation work (particularly the heavy  
 106 use of Shapley values).

## 107 2.3 Algorithmic Collective Action

108 Algorithmic Collective Action (ACA) is defined as efforts by groups to achieve common goals  
 109 by strategically manipulating their data contributions to change algorithmic outputs [26, 12]. This  
 110 approach builds on the “data leverage” concept, where individuals treat their data not as passive  
 111 inputs but as tools for influencing algorithmic outcomes [27].

112 Research shows that coordination in ACA is essential for effectiveness. While individual users lack  
 113 sufficient leverage to influence algorithmic behavior due to the high volume of data, even a small  
 114 collective can influence outcomes through coordinated action [12, 1]. Fan communities also provide  
 115 a compelling example of large-scale algorithmic collective action, where millions of users coordinate  
 116 to manipulate social media algorithms and boost visibility for their preferred content [30].

117 ACA literature identifies several other challenges that affect its ability to influence algorithmic  
 118 systems. One is informational: participants must act with limited reliable knowledge about how the  
 119 opaque and rapidly shifting algorithm works [30, 12]. Without formal tools or feedback mechanisms,  
 120 groups develop speculative explanations or informal rules of thumb, which may be inconsistent or  
 121 short-lived. Similarly, organizing collective behavior presents difficulties as recruiting large numbers  
 122 requires significant effort to align motivations, coordinate timing, and sustain engagement. Success  
 123 also depends on organizers’ ability to translate complex algorithmic behaviors into clear, actionable  
 124 tutorials for general users [30].

### 3 Position Statement

**Algorithmic Collective Action and Explainable AI both share common goals: advancing human understanding, strengthening user agency, and promoting transparency and accountability in algorithmic systems. Their complementary perspectives (action vs explanation) can offer a more effective path towards achieving these objectives.**

The conceptual overlap between these communities centers on understanding influence. XAI tends to focus on the “why” behind algorithmic decisions and explains how different factors may influence outcomes [6, 13]. ACA emphasizes the “how” of system change and develops specific strategies for influencing algorithmic behaviors [30, 12]. These two perspectives complement each other, and by combining their work, XAI can benefit from ACA’s emphasis on practical interventions that reveal which explanations work better for real-world organizing, while ACA can leverage XAI work to design more informed and strategic interventions.

Ultimately, these communities both advance a common goal of empowering people, especially groups with limited individual influence, to meaningfully understand and shape algorithmic systems.

We also emphasize the importance of collaboration across industry, developers, researchers, practitioners, and user collectives to align platform goals with user-led efforts to address algorithmic bias and fairness. Through these partnerships, collectives can pursue equity goals more effectively, which benefits both users and industry platform owners.

### 4 Common Ground

#### 4.1 Goals

We identify two goals shared by both communities:

**Goal 1: Enhancing User Agency and Control Over Algorithmic Systems.** Both fields help transform users from passive recipients of algorithmic decisions into active agents capable of understanding and influencing these systems. XAI pursues this goal by providing “human-understandable justification for a system’s behavior” [7]. Data valuation contributes by showing users how small modifications, such as reweighting or removing specific points, would alter model parameters and predictions [19]. ACA research highlights the growing public dissatisfaction with the lack of agency in algorithmic decision-making and explores how users engage in forms of algorithmic resistance or activism to reclaim influence.

**Goal 2: Promoting Algorithmic Transparency and Accountability.** While Goal 1 centers on empowering users to act within algorithmic systems, Goal 2 focuses on questioning and reforming these systems. XAI enables scrutiny of outputs for bias, fairness, and safety concerns through explanations [6]. ACA complements this by showing how users, through collective activism, challenge the legitimacy and integrity of these systems, for example, by revealing gaps in platform governance and participating in strategic data manipulation campaigns [30].

#### 4.2 Challenges

In addition to these shared goals, XAI and ACA also face similar challenges, especially around measuring success and developing useful mental models of algorithmic systems. ACA participants often express frustration over questions; past work shared quotes such as “Is it really useful for me to do this data?” and “Is the time and effort I spent really rewarding my idol accordingly?” [30]. Here, individual actors lack information about the broader picture. In XAI, researchers face challenges in evaluating whether explanations genuinely enhance user understanding and decision-making [25], which serve as indicators of XAI success. There is also an added complication that explanation methods may not accurately represent the true underlying model [22].

Furthermore, both XAI and ACA benefit when researchers assess mental models of complex algorithmic systems formed by users. For example, prior XAI work [20] found that although users may feel that explanations are helpful and easy to understand, explanations may not actually help users develop more accurate mental models. ACA participants face the same challenge as demonstrated in previous work [30], where fans develop their understanding through iterative observation and “sensemaking loops”. However, algorithmic opacity means they often rely on folk theories and trial-and-error.

### 4.3 Concrete Project Examples

Here, we present examples of XAI and ACA projects, highlighting their differences and opportunities for collaboration.

**XAI projects** often explain individual predictions to individual users. For example, in cancer detection models, the explanation could highlight which input pixels most contribute to the model’s prediction [6], while credit scoring applications may generate counterfactual explanations that specify required income thresholds for loan approvals [24]. XAI projects also typically include measures for assessing explanation quality, such as instruments designed to measure whether explanations actually enable user understanding or action [24].

**ACA projects** focus on coordinated group efforts to influence algorithmic predictions rather than on understanding these systems. For example, in recommender systems, collectives may coordinate by agreeing on target songs to strategically place in playlists [2], or as an adversarial example, data poisoning campaigns where groups may inject fake users with crafted ratings to manipulate recommender systems towards promoting target items to many users [14].

**These projects, we argue, differ in three ways:** their information flow, scale, and success metrics. XAI projects typically operate at inference time, sending information from the system to the user (i.e., explaining model decisions or suggested actions in algorithmic recourse). Notably, counterfactual explanations might appear to support this reverse flow, since they identify changes a user could make to achieve a different outcome; they do not necessarily provide feasible guidance on how to implement these changes, and thus offer limited practical agency [24]. Directive explanations, an extension to counterfactuals, attempt to address this by offering concrete steps the individual could perform, but their feasibility and real-world effectiveness remain underexplored [24]. Even when XAI methods gesture towards actionability, information flow remains largely one-way, from system to user. In comparison, ACA tends to intervene at training time, flowing influence from the user to the system (i.e., changing decisions through coordinated changes to the training data).

In terms of scale, XAI projects typically serve individuals and measure success through user understanding, mental model accuracy, or model improvement (in the case of debugging), while ACA projects typically require significant coordination (i.e., a larger group of users) since individual contributions represent only an infinitesimal fraction of the platform data and will measure success through perceived algorithmic outcome changes.

**Combining XAI’s explanatory strength with ACA’s coordination focus** reveals opportunities for better meeting the goals we described above. For instance, current fan collectives operate with a limited understanding, relying on trial and error methods, suggesting XAI techniques could help organizers better understand algorithms before developing their collective’s strategy [30]. It is worth noting that different types of explanations may vary in their usefulness for ACA. Simpler forms of explanations may be more useful for broad groups who need a shared understanding to coordinate, whereas more technical or detailed explanations may be better suited to entirely different settings with specialized or domain specific contexts, such as a developer debugging a model or an expert making high-stakes decisions in medical imaging.

As another example, data value estimation techniques, such as influence functions, could be used to help collectives understand which training examples their actions could affect in simplified models. Similarly, in the playlist manipulation strategy, XAI could explain why certain placement strategies work, while algorithmic recourse could extend beyond individual explanations to collective scenarios.

## 5 Alternate Views and Counterarguments

This section addresses “counterarguments” to the stated positions, i.e., reasons why some might view these two fields as relatively distinct and not necessarily well-suited for connecting work.

**Counterargument 1: Different Target Users and Use Cases.** Critics might argue that XAI and ACA serve fundamentally different user groups—XAI primarily focuses on individuals seeking to understand algorithmic outputs, while ACA focuses on coordinating collective efforts to influence algorithms. Because individual explanatory needs can differ from requirements for informing group action, these critics may argue that the two fields have limited overlap or mutual relevance.

226 **Our Response:** While the immediate users may vary, both XAI and ACA ultimately contribute toward  
227 the shared objective of algorithmic accountability and user empowerment. Importantly, successful  
228 collective action depends on some level of individual understanding. For example, research on  
229 fan communities shows how participants build personal ideas of algorithmic behaviors to engage  
230 in coordinated algorithmic intervention [30]. Rather than being at odds, individual and collective  
231 needs form a complementary progression, where personal understanding may empower meaningful  
232 participation in collective efforts.

233 **Counterargument 2: Concerns About Manipulation.** A concern for industry platform owners  
234 is that integrating XAI and ACA could facilitate algorithmic manipulation for harmful purposes.  
235 Explanatory tools might enable more sophisticated and coordinated attacks on digital platforms using  
236 ML systems, creating security and operational risks.

237 **Our Response:** This concern applies to existing XAI research, which already provides insights that  
238 could be misused. For example, this concern may be valid in contexts where highly coordinated  
239 malicious actors, such as nation-states, may exploit explanatory tools to exert control. However,  
240 this view also overlooks systemic issues at hand. ACA requires some critical mass of users to be  
241 effective, and if a significant number of benign users must resort to a coordinated attack to make their  
242 voices heard, this signals a failure of the platform’s existing feedback and governance mechanisms.  
243 The solution may not be to withhold transparency tools out of a fear of misuse, but to develop XAI  
244 solutions that empower collective action while incorporating safeguards against misuse, recognizing  
245 that both concerns and benefits will coexist [30].

## 246 6 Recommendations

247 This integration requires coordinated development of new tools, research methodologies, and practical  
248 applications across multiple stakeholder groups:

249 **Responsible AI teams in Industry** stand to benefit from recognizing that current collective action  
250 efforts emerge without proper guidance, potentially causing unintended consequences that may harm  
251 both users and platforms. While some coordinated collective action campaigns may be designed to  
252 produce outcomes that are undesirable for organizations operating AI systems, many campaigns may  
253 fundamentally aim to address algorithmic inequities. These goals directly align with responsible AI  
254 initiatives, suggesting that industry teams could engage to help produce more constructive outcomes  
255 while preventing potential misuse. However, it is worth noting that such engagement may encounter  
256 challenges including platform level countermeasures, organizational resistance to open collaboration,  
257 and potentially limited transparency policies. For maintaining ongoing collaboration, teams should  
258 work towards gathering or defining clear organizational guidelines and operational protocols that  
259 promote transparent and responsible collaboration while protecting sensitive proprietary information.

260 **Developers** interested in either field should prototype tools that enable both transparency and  
261 coordination. We recommend: creating dashboards that both explain algorithmic output and also  
262 suggest viable collective action paths (e.g., how many coordinated playlist edits are needed to flip  
263 the prediction), building open-source toolkits that allow groups to adapt explanations and organizing  
264 tools to their context, and easily modifiable interfaces and visualizations for their needs. Notably,  
265 developers may also encounter challenges such as resource constraints, restricted data access, and  
266 dependencies on closed APIs that can limit experimentation. Building partnerships with platforms or  
267 using synthetic or federated datasets could help mitigate these barriers.

268 **Researchers** in both XAI and ACA communities should collaboratively develop shared methodolo-  
269 gies and tools. Specific research priorities could include: creating simulation environments where  
270 collectives can test algorithmic strategies using XAI feedback [16, 26, 28, 4], studying how different  
271 explanation modalities (visual, textual, interactive) [29] support coordination among participants,  
272 and designing longitudinal studies that track how XAI interfaces and tools help with collective  
273 action. Conducting such studies, however, presents ethical and logistical challenges and researchers  
274 must carefully address risks of perpetuating existing biases, and ensure equitable access to tools  
275 and participation. Dedicated ethical review processes and controlled experimental environments are  
276 needed to minimize negative social impacts and preserve ecological validity.

277 **Practitioners** involved in organizing ACA can bridge research and development by providing  
278 developers with community feedback on toolkits, interfaces, and visualizations. They can also

inform researchers about practical needs and challenges, ensuring XAI tools meaningfully and materially support collective goals. As discussed, organizers in practice may struggle with participant coordination, sustaining engagement, or scaling tools across different cultural and regulatory settings. Participation frameworks and governance models will play a central role in translating experimental tools into real world collective actions.

## 7 Future Work

An important direction for future research involves developing structured workflows that clearly specify the practical steps and interactions needed to connect XAI and ACA, creating an operational bridge between the disciplines. Additionally, design frameworks can help provide guidelines for building systems that integrate explanation and coordination in conjunction. Another promising path for implementing these ideas is leveraging the advanced language understanding capabilities of language models, especially when implemented directly into LLM-based platforms, to facilitate explanation, coordination and interactive support for users. Lastly, empirical evaluations through user studies and simulations are needed to evaluate the feasibility and effectiveness of these frameworks, surface unforeseen challenges, and guide iterative improvements.

## 8 Conclusion

ACA and XAI research address complementary aspects of the same fundamental challenge: empowering users to understand and influence algorithmic systems. Our paper identifies the conceptual overlaps between these fields and proposes actionable examples for integration, showing how XAI’s explanatory capabilities can enhance ACA’s effectiveness, while ACA’s practical experiences can inform more actionable XAI design.

We find that both fields may already be moving towards this convergence through work on algorithmic recourse and argue that they ultimately share similar goals and face many of the same challenges. Addressing them involves further work on integrating explanation and action. Creating deeper connections between XAI and ACA through shared tools, collaborative research, and practical applications can help build systems that not only reveal how algorithms work but also empower people to change them.

## 9 Acknowledgments

We thank the reviewers for their valuable feedback and suggestions, which have helped improve this workshop paper.

## References

- [1] Joachim Baumann and Celestine Mendler-Dünnér. 2024. Algorithmic collective action in recommender systems: promoting songs by reordering playlists. *Advances in Neural Information Processing Systems* 37 (2024), 119123–119149.
- [2] Joachim Baumann and Celestine Mendler-Dünnér. 2025. Algorithmic Collective Action in Recommender Systems: Promoting Songs by Reordering Playlists. *arXiv:2404.04269 [cs]* doi:10.48550/arXiv.2404.04269
- [3] Aditya Bhattacharya. 2024. Towards directive explanations: Crafting explainable ai systems for actionable human-ai interactions. In *Extended Abstracts of the CHI Conference on Human Factors in Computing Systems*. 1–6.
- [4] Dangxing Chen, Jingfeng Chen, and Weicheng Ye. 2024. Why Groups Matter: Necessity of Group Structures in Attributions. *arXiv:2408.05701 [q-fin]* doi:10.48550/arXiv.2408.05701
- [5] Elliot Creager and Richard Zemel. 2023. Online algorithmic recourse by collective action. *arXiv preprint arXiv:2401.00055* (2023).

- [6] Arun Das and Paul Rad. 2020. Opportunities and Challenges in Explainable Artificial Intelligence (XAI): A Survey. *arXiv:2006.11371 [cs]* doi:10.48550/arXiv.2006.11371
- [7] Upol Ehsan, Koustuv Saha, Munmun De Choudhury, and Mark O. Riedl. 2023. Charting the Sociotechnical Gap in Explainable AI: A Framework to Address the Gap in XAI. 7 (2023), 1–32. Issue CSCW1. *arXiv:2302.00799 [cs]* doi:10.1145/3579467
- [8] João Fonseca, Andrew Bell, Carlo Abrate, Francesco Bonchi, and Julia Stoyanovich. 2023. Setting the right expectations: Algorithmic recourse over time. In *Proceedings of the 3rd ACM Conference on Equity and Access in Algorithms, Mechanisms, and Optimization*. 1–11.
- [9] Krishna Gade, Sahin Cem Geyik, Krishnamurthy Kenthapadi, Varun Mithal, and Ankur Taly. 2019. Explainable AI in industry. In *Proceedings of the 25th ACM SIGKDD international conference on knowledge discovery & data mining*. 3203–3204.
- [10] Amirata Ghorbani, Michael Kim, and James Zou. 2020. A distributional framework for data valuation. In *International Conference on Machine Learning*. PMLR, 3535–3544.
- [11] Vivek Gupta, Pegah Nokhiz, Chitradheep Dutta Roy, and Suresh Venkatasubramanian. 2019. Equalizing recourse across groups. *arXiv preprint arXiv:1909.03166* (2019).
- [12] Moritz Hardt, Eric Mazumdar, Celestine Mendler-Dünner, and Tijana Zrnica. 2024. Algorithmic Collective Action in Machine Learning. *arXiv:2302.04262 [cs]* doi:10.48550/arXiv.2302.04262
- [13] Zhongli Filippo Hu, Tsvi Kuflik, Ionela Georgiana Mocanu, Shabanam Najafian, and Avital Shulner Tal. 2021. Recent Studies of XAI - Review. In *Adjunct Proceedings of the 29th ACM Conference on User Modeling, Adaptation and Personalization* (New York, NY, USA, 2021-06-22) (*UMAP '21*). Association for Computing Machinery, 421–431. doi:10.1145/3450614.3463354
- [14] Hai Huang, Jiaming Mu, Neil Zhenqiang Gong, Qi Li, Bin Liu, and Mingwei Xu. 2021. Data poisoning attacks to deep learning based recommender systems. *arXiv preprint arXiv:2101.02644* (2021).
- [15] Matthew Jagielski, Giorgio Severi, Niklas Pousette Harger, and Alina Oprea. 2021. Subpopulation Data Poisoning Attacks. *arXiv:2006.14026 [cs]* doi:10.48550/arXiv.2006.14026
- [16] Aditya Karan, Nicholas Vincent, Karrie Karahalios, and Hari Sundaram. 2025. Algorithmic Collective Action with Two Collectives. In *Proceedings of the 2025 ACM Conference on Fairness, Accountability, and Transparency* (2025-06-23). 1468–1483. *arXiv:2505.00195 [cs]* doi:10.1145/3715275.3732098
- [17] Amir-Hossein Karimi, Gilles Barthe, Bernhard Schölkopf, and Isabel Valera. 2022. A survey of algorithmic recourse: contrastive explanations and consequential recommendations. *Comput. Surveys* 55, 5 (2022), 1–29.
- [18] Amir-Hossein Karimi, Gilles Barthe, Bernhard Schölkopf, and Isabel Valera. 2021. A survey of algorithmic recourse: definitions, formulations, solutions, and prospects. *arXiv:2010.04050 [cs]* doi:10.48550/arXiv.2010.04050
- [19] Pang Wei Koh and Percy Liang. 2020. Understanding Black-box Predictions via Influence Functions. *arXiv:1703.04730 [stat]* doi:10.48550/arXiv.1703.04730
- [20] Retno Larasati, Anna De Liddo, and Enrico Motta. 2023. Meaningful Explanation Effect on User’s Trust in an AI Medical System: Designing Explanations for Non-Expert Users. 13, 4 (2023), 30:1–30:39. doi:10.1145/3631614
- [21] Q. Vera Liao and Kush R. Varshney. 2022. Human-Centered Explainable AI (XAI): From Algorithms to User Experiences. *arXiv:2110.10790 [cs]* doi:10.48550/arXiv.2110.10790
- [22] Dino Pedreschi, Fosca Giannotti, Riccardo Guidotti, Anna Monreale, Salvatore Ruggieri, and Franco Turini. 2019. Meaningful explanations of black box AI decision systems. In *Proceedings of the AAAI conference on artificial intelligence*, Vol. 33. 9780–9784.

- 372 [23] Kaivalya Rawal and Himabindu Lakkaraju. 2020. Beyond individualized recourse: Interpretable  
373 and interactive summaries of actionable recourses. *Advances in Neural Information Processing*  
374 *Systems* 33 (2020), 12187–12198.
- 375 [24] Ronal Singh, Tim Miller, Liz Sonenberg, Eduardo Velloso, Frank Vetere, and Piers Howe. 2025.  
376 An Actionability Assessment Tool for Enhancing Algorithmic Recourse in Explainable AI.  
377 *IEEE Transactions on Human-Machine Systems* (2025).
- 378 [25] Ronal Singh, Tim Miller, Liz Sonenberg, Eduardo Velloso, Frank Vetere, Piers Howe, and Paul  
379 Dourish. 2024. An Actionability Assessment Tool for Explainable AI. arXiv:2407.09516 [cs]  
380 doi:10.48550/arXiv.2407.09516
- 381 [26] Rushabh Solanki, Meghana Bhange, Ulrich Aïvodji, and Elliot Creager. 2025. Crowding Out  
382 The Noise: Algorithmic Collective Action Under Differential Privacy. arXiv:2505.05707 [cs]  
383 doi:10.48550/arXiv.2505.05707
- 384 [27] Nicholas Vincent, Hanlin Li, Nicole Tilly, Stevie Chancellor, and Brent Hecht. 2021. Data  
385 Leverage: A Framework for Empowering the Public in its Relationship with Technology  
386 Companies. arXiv:2012.09995 [cs] doi:10.48550/arXiv.2012.09995
- 387 [28] Greta Warren, Mark T. Keane, Christophe Gueret, and Eoin Delaney. 2023. Explaining Groups  
388 of Instances Counterfactually for XAI: A Use Case, Algorithm and User Study for Group-  
389 Counterfactuals. arXiv:2303.09297 [cs] doi:10.48550/arXiv.2303.09297
- 390 [29] Katharina Weitz, Dominik Schiller, Ruben Schlagowski, Tobias Huber, and Elisabeth André.  
391 2021. “Let me explain!”: exploring the potential of virtual agents in explainable AI interaction  
392 design. *Journal on Multimodal User Interfaces* 15, 2 (2021), 87–98.
- 393 [30] Qing Xiao, Yuhang Zheng, Xianzhe Fan, Bingbing Zhang, and Zhicong Lu. 2025. Let’s  
394 Influence Algorithms Together: How Millions of Fans Build Collective Understanding of  
395 Algorithms and Organize Coordinated Algorithmic Actions. arXiv:2409.10670 [cs] doi:10.  
396 48550/arXiv.2409.10670
- 397 [31] Tong Zhang, Mengao Zhang, Wei Yan Low, X. Jessie Yang, and Boyang Li. 2025. Conversa-  
398 tional Explanations: Discussing Explainable AI with Non-AI Experts. arXiv:2503.16444 [cs]  
399 doi:10.48550/arXiv.2503.16444