

Towards Foundation-model-based Multiagent System to Accelerate AI for Social Impact

Blue Sky Ideas Track

Yunfan Zhao

Harvard University, GE Healthcare
yunfanzhao@fas.harvard.edu

Aparna Taneja

Google Deepmind
aparnataneja@google.com

Niclas Boehmer

Harvard University, Hasso Plattner Institute
niclas.boehmer@hpi.de

Milind Tambe

Harvard University, Google Deepmind
milind_tambe@harvard.edu

ABSTRACT

AI for social impact (AI4SI) offers significant potential for addressing complex societal challenges in areas such as public health, agriculture, education, conservation, and public safety. However, existing AI4SI research is often labor-intensive and resource-demanding, limiting its accessibility and scalability; the standard approach is to design a (base-level) system tailored to a specific AI4SI problem. We propose the development of a novel meta-level multi-agent system designed to accelerate the development of such base-level systems, thereby reducing the computational cost and the burden on social impact domain experts and AI researchers. Leveraging advancements in foundation models and large language models, our proposed approach focuses on resource allocation problems providing help across the full AI4SI pipeline from problem formulation over solution design to impact evaluation. We highlight the ethical considerations and challenges inherent in deploying such systems and emphasize the importance of a human-in-the-loop approach to ensure the responsible and effective application of AI systems.

KEYWORDS

Social impact, Foundation Models, Multiagent Systems

ACM Reference Format:

Yunfan Zhao, Niclas Boehmer, Aparna Taneja, and Milind Tambe. 2025. Towards Foundation-model-based Multiagent System to Accelerate AI for Social Impact: Blue Sky Ideas Track. In *Proc. of the 24th International Conference on Autonomous Agents and Multiagent Systems (AAMAS 2025)*, Detroit, Michigan, USA, May 19 – 23, 2025, IFAAMAS, 7 pages.

1 INTRODUCTION

Artificial intelligence (AI) for social impact (AI4SI), which focuses on leveraging AI to address societal issues, has gained traction in both academia and industry [9, 14, 26, 37, 57]. With advancements in AI and multi-agent systems, there is an opportunity to apply these technologies to complex problems in areas like public safety, wildlife conservation, and public health [25, 36, 42, 72]. Previously,

AI4SI research has been very labor-intensive, as it is oftentimes necessary to develop customized approaches going beyond conventional methods to address the challenges characteristic to these domains such as low resources and noisy or scarce data. This limits the overall impact of AI4SI research, as every individual effort requires non-trivial time, expertise, and financial investment. We envision the formation of a new approach to AI4SI which is less labor-intensive, customizable by non-experts, and can thus be made more widely available. We believe promising progress can be made on this vision by employing recent methodological advancements in computer science research, tapping into a substantial currently unleveraged potential.

In this paper, we will use resource allocation problems, which often arise in AI4SI domains [4, 5, 16, 41, 52, 57, 72] as our running example, yet our general ideas also apply more broadly. Some examples of previously studied resource allocation problems in AI4SI include strategically scheduling patrols in protected conservation areas [27, 29, 31, 88] and distributing scarce healthcare resources to optimize people’s health outcomes [50, 63, 70, 96].

We envision a meta-level multi-agent system that helps us accelerate the development of base-level systems, which tackle specific AI4SI problems. The meta-level system would help non-profits and AI researchers in social impact domains leverage AI without having to invest significant amounts of labor and resources to build a tailored base-level system from scratch. Our envisioned system leverages foundation models, which are typically developed by pretraining on available datasets, and can be used on different downstream tasks or new challenges [8, 39, 71, 97]. Our proposed system involves: (i) using LLM based meta-level agents to communicate with decision makers in human language to understand the problem from the perspective of decision-makers; (ii) employing meta-level agents and foundation models for AI4SI problems to design base-level systems for AI4SI problems; and (iii) using meta-level agents for field-testing solutions to validate their impact.

Importantly, instead of taking over the AI4SI pipeline, the meta-level system will accelerate the process, improve generalization, and enable thorough evaluation, with human-in-the-loop required in each part of the system. We will discuss and highlight major challenges and our vision for each phase of this system, emphasizing the multi-agent aspect. We also discuss ethics and fairness aspects.



This work is licensed under a Creative Commons Attribution International 4.0 License.

The current era of foundation models has led to contemplation on the challenges and opportunities that such models may offer in multiple areas, from medical AI [51] to autonomous supply chains [92], autonomous mining [45], and robotics [74]. In comparison with these previous works, this paper focuses on the use of foundation models in AI4SI. Moreover, even within AI4SI, we exemplarily focus on optimizing the allocation of limited resources, providing an analysis of challenges in a specific aspect of AI4SI. Furthermore, in contrast to previous work, we focus more on foundation-model-based agents at the meta-level, to assist base-level systems that optimize these limited resources, rather than replacing the base-level system entirely. This enables existing well-developed resource optimization tools to be brought to bear on relevant challenges as required, allowing instead the foundation models to configure the tools as needed.

2 PRELIMINARIES

We formally define the key concepts of meta-level and base-level systems, which we will refer to throughout this work.

DEFINITION 1 (META-LEVEL AND BASE-LEVEL SYSTEMS). *A base-level AI4SI system is the actual deployed system that will solve the problem on the ground. A meta-level system helps us accelerate the development of a base-level system or accelerate its modification as needed for a new objective.*

Notably, a meta-level system does not actually solve an AI4SI problem. A meta-level system may involve several meta-level agents, each responsible for different tasks in the development of the base-level system, which may interact with each other. In the context of the use cases we are focusing on in this paper, the base system itself will be a multi-agent system or one that models multi-agent interactions such as a social network. The base-level agents present in the base-level system are defined as follows:

DEFINITION 2 (BASE-LEVEL AGENTS). *In the base-level system, the base-level agents model or serve as abstract representations of individuals or entities in the real-world.*

Whereas the idea of a meta-level architecture has been proposed in agent architectures, there the idea is to directly assist agents in their immediate problem solving [13, 28, 59, 78]. In our work, the meta-level refers to deciding which agents and how many to select, how to evaluate their performance, and other such tasks. Another difference is that at least as conceived now, our meta-agents are focused on assisting humans in building base-level agents.

Another key concept we will frequently refer to is foundation-model-based agents (FM agents) used in the meta-level system.

DEFINITION 3 (FOUNDATION-MODEL-BASED AGENTS (FM-AGENTS)). *In the meta-level system, we use meta-level agents that employ foundation models including LLMs. We will refer to these agents as FM-agents*

An AI4SI pipeline typically involves three phases [9, 57]:

Formulate the problem Identify how to best represent or model the real-world entities and multiple stakeholders involved in the on-the-ground problem, along with their constraints and objectives. For example, the real-world entities may include individuals enrolled in social welfare programs and decision-makers such as non-profit program managers. Previously,

this effort was largely manual and involved multiple parties, including nonprofits and AI researchers [19, 64, 68].

Design solution method Identify appropriate solution methods and adapt them to design the base-level system. Previously, this process was largely manual in terms of researcher efforts to tailor advances in AI algorithms to the specific application scenarios. The adaptation to different application scenarios was also largely manual [20, 54, 75, 81, 87].

Evaluation, refinement, and deployment Testing, improving, and deploying solutions has typically been done manually, adding to the workload of researchers and even more importantly to that of resource-constraint NGO workers [72].

Notably, each step in the AI4SI pipeline poses its own challenges [66], implying that AI4SI work often requires substantially more effort than pure AI algorithmic improvement research. Specifically, formulating the problem and collecting detailed information (e.g. potential decisions available and their impact on each of the individuals) can be time-consuming and expensive [60]. Moreover, designing solution methods can require significant time from AI experts and social impact domain experts [68], who must work together to devise a tailored solution method for every application scenario. Thorough testing and evaluation before deployment often require substantial manual effort, as detailed simulation studies are often required (e.g. by regulations or as precautionary measures) [6], and AI experts typically manually design each simulation study from scratch.

For our running example, we consider allocating limited interventions (specifically live service calls made by health workers) in ARMMAN, which is a non-profit in India focusing on improving health awareness for expectant and new mothers [62, 97]. Their health workers make service calls to boost the engagement of mothers enrolled in their health information program. They have shown that AI powered solutions can reduce engagement drops by about 30% in real-world deployment [49, 75].

3 FORMULATING THE PROBLEM

In this phase, we discuss how to formulate an AI4SI problem. Existing works require researchers and human experts in AI to talk to collaborating non-profit organizations to understand who makes the decisions within a problem and gather key information on the agents such as demographics [72, 80]. This process is expensive and labor-intensive, and non-profits may not have AI-trained staff to assist with this, making it difficult to ensure that AI solutions are accurately tailored to the complexities of the real-world problem. Thus, previous works often require AI researchers to have numerous rounds of discussions with non-profits and arrange a field trip to speak with key stakeholders in social good programs. Whereas these discussions are fundamental to AI4SI, some of the work oriented toward formulating the right base-level model is repetitive. To accelerate this work we propose the following vision:

VISION 1. *Employ FM-agents to (i) identify the base-level agents involved, (ii) find an adequate formal description for the setting (e.g. as a Markov Decision Process), and (iii) define key components of the settings (e.g. state space, action space, and reward function in a MDP).*

The FM-agents may use large language models to communicate with a partnering non-profit to formulate the social challenge as

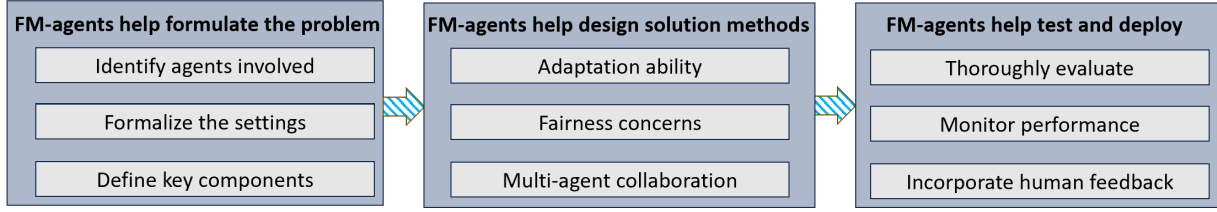


Figure 1: Overview of our proposed AI for social impact (AI4SI) workflow. The three key phases, formulating the problem, designing the solution methods, and testing and deployment, are discussed in Sections 3, 4, and 5 respectively.

an AI problem and to understand who makes the decisions within a problem. The world knowledge of LLMs may help the FM-agent uncover confounding variables and undocumented information [33–35, 43, 73, 91]. The FM-agent should determine what information the individuals or entities in the real-world have to guide their decisions. The FM-agent should then define base-level agents to model the individuals or entities and make proper assumptions about the information available to them.

The FM-agent may gather data from past interactions, including the effects of actions, and observed costs or rewards, and interactions between agents [24, 58, 98]. The FM-agent, potentially LLM-based, could communicate directly with beneficiaries enrolled in non-profit’s program in their native language to gain a clearer perspective and would not have time constraints.

RUNNING EXAMPLE 1. *In ARMMAN, the base-level agents represent beneficiaries enrolled, and the state and action space correspond to beneficiaries’ engagement levels and possible schedules of service calls, respectively. The FM-agent should have conversations with domain experts in ARMMAN to find out that one way to approach this application scenario is to model beneficiaries as agents that follow Markov Decision Processes. After that, the FM-agent should define the state space, action space, and other key parts of the MDP.*

4 DESIGNING SOLUTION METHODS

In this phase, we elaborate on how to design a solution method for AI4SI problems. Existing approaches often require manually designing solution methods tailored to each application scenario [2, 46, 47, 82, 86, 93]. This approach fails to easily adapt to new application scenarios or knowledge and data from previous application scenarios, motivating the development of foundation models to accelerate solution approaches for problems in AI4SI. **Besides adaptation ability, other aspects that are of particular importance when designing solution methods in AI4SI problems include ethics, fairness, and collaboration among base-level agents. We will motivate each of these aspects and propose our visions to address them.**

Adaptation is crucial in AI for social impact domains due to the dynamic and evolving nature of these environments. Social issues are often complex and multifaceted, with priorities evolving over time [6, 7, 15]. A foundation model for resource allocation could accelerate developing solutions for different application scenarios without incurring repeated development costs, making them more affordable and accessible to low-resource communities [8, 97]. For example, a foundation model designed to analyze medical data

can be adapted to different diseases, health conditions, or populations, improving health outcomes on a larger scale [44]. Based on advances in foundation models and adaptation, we propose the following vision:

VISION 2. *Build a foundation model for resource allocation problems in AI4SI domains that can be adapted to and finetuned on specific application scenarios. For each new AI4SI application scenario, employ a FM-agent to leverage the foundation model and provide solution methods.*

RUNNING EXAMPLE 2. *A concrete example of foundation model for resource allocations tasks is given by Zhao et al. [96], who develop a pretrained restless bandit model that can be finetuned on various resource allocation application scenarios that ARMMAN may encounter. Currently, the application scenarios have different number of base-level agents and different amounts of distribution shifts. Here, our research idea is to start with such a foundation model and allow an FM-agent to adapt and specialize it to newer scenarios that may involve bigger changes than just differences in numbers of base-level agents. This could include new application scenarios that may need a change of the states and actions in the restless bandit model.*

To address the inherent complexity of social challenges, we may also use FM-agents in the form of Large Language Model (LLM) to process and understand human instructions, queries, and feedback from stakeholders to alter the priorities within the resource allocation process [67, 79, 83, 90]. For example, in the ARMMAN domain, a program manager may suggest prioritizing a specific segment of the underserved population such as those older in age, which an LLM could interpret and accordingly adjust the restless bandit resource allocation model by changing its reward function [6, 69].

Besides the adaptation aspect, fostering efficient collaboration among multiple base-level agents plays an important role in AI4SI research. Recall that a base-level agent serves as an abstract representation of an individual or an entity in the real-world. In some problems, multiple base-level agents may collectively plan to counter an adversary such as wildlife poacher or terrorists [31, 65, 85]. In other problems, multiple base-level agents may communicate to mitigate the impact of data errors, which frequently arise in real-world situations due to factors such as inconsistent data collection methods and deliberate noise introduced for differential privacy [17, 18, 21, 22, 55, 99].

However, previous works in AI4SI often require human experts in AI to manually craft ways of collaboration among base-level agents for specific AI4SI domains. An FM-agent can help with this process to accelerate AI for social impact work:

VISION 3. *The FM-agent, when designing solution methods for base-level systems, should design effective communication channels and strategies for base-level agents. Specifically, these channels should allow base-level agents to learn from each other's experience and to improve decision-making.*

Although the above vision on developing collaboration may appear to be straightforward for human-AI experts crafting solution methods, it is not easy for FM-agents to figure out due to complex relationships between base-level agents [38, 76, 77]. The FM-agent may use the world knowledge of LLMs to understand which communications can be potentially useful and should be included in the design of solution methods [10]. Here a communication channel may be that an individual or entity talks to another via cellular network or other infrastructure in place.

4.1 Ethics and Fairness

In high-stakes resource allocation scenarios like healthcare, authorities frequently prioritize certain groups based on sensitive attributes, aiming to address the needs of those most disadvantaged [1, 70]. For example, governments may mandate non-discrimination based on sensitive attributes, while non-profits may prioritize low-income groups. Given the importance of fairness in base-level system design and the solution method's tangible impact on people's lives, we propose the following idea:

VISION 4. *Ensure that the FM-agent recommends the design of a base-level system that does not discriminate against any subpopulation or result in unfavorable outcomes for under-privileged groups.*

When accelerating the design of base-level systems, the FM-agent should ensure fairness guarantees or fairness checks are in place. This can be done by explicitly incorporating fairness in designing base-level systems for social impact applications [70, 94]. However, this added complexity is difficult for FM-agents to handle, due to the fact that AI may not easily understand demographic information available in text or abstract fairness concepts potentially based on demographics. Blindly applying fairness constraints or solely optimizing a fairness objective may greatly compromise the overall effectiveness of the solution method. Thus, fairness concerns necessitate innovative strategies to ensure that FM-agents design base-level systems that balance fairness and utility.

RUNNING EXAMPLE 3. *In ARMMAN a concrete example of fairness constraints is the enforcement of non-discrimination based on sensitive attributes and the prioritization of low-income and low-education groups to reduce socio-economic disparities. In the ARMMAN context, demographic information for enrolled beneficiaries are available. When designing a base-level system, the FM-agent could enforce fairness constraints such as that each beneficiary must receive a sufficient amount of resources within some time. Additionally, the FM-agent should prioritize underprivileged groups by explicitly optimizing a fairness objective (e.g. Max Nash Welfare or Maximin Reward) in the solution method.*

5 TESTING AND DEPLOYMENT

In this phase, we explain how to thoroughly evaluate and deploy AI models for social impact domains. We use FM-agents to improve model testing and facilitate real-world deployment.

Deploying an AI model in real-world social impact domains without sufficient simulation studies may result in poor decision-making on crucial public resources. Thus, thoroughly testing and evaluating AI algorithms or trained models is an important aspect of accelerating AI for social impact. Based on above, we propose the following research idea:

VISION 5. *Employ FM-agents, based on LLMs, to simulate agents' behaviors. Here we wish to build a powerful simulator that serves as a good proxy of real-world deployment environment and can effectively evaluate the performance of trained models.*

LLM based simulators have recently received great interest, and there is demonstrated success in using LLMs to simulate human behaviors in fields including education, healthcare, and social sciences [3, 11, 48, 61, 89]. To build such a LLM simulator / evaluator for AI4SI problems, we should represent observations and possible decisions in a way that the LLM can understand, and potentially use textual descriptions combined with structured data, to help the LLM simulator generate contextually appropriate (e.g. suitable for the domain) individual behaviors. Furthermore, textual descriptions on individual's characteristics, such as demographic information (age, gender, geographical location, etc) may help LLMs to better understand how individuals' condition would evolve over time. For a particular AI4SI problem, we may need to finetune LLMs on historical data collected to better simulate the individual's trajectories.

RUNNING EXAMPLE 4. *In the ARMMAN application, we need to thoroughly test algorithms before real-world deployment. An FM-agent can employ LLMs to perform agent-based simulations to evaluate learning algorithms [32, 53], or use cognitive models to augment ML based evaluation approaches [30, 62].*

Having discussed evaluation before real-world deployment, we now move on to challenges in the deployment. During the deployment of AI models, there can be shifts in the user base or shifts in people's behaviors [23]. Different from adaptation ability taken into account during model development and training, distribution shifts in testing can be unexpected and the need to handle these shifts can be urgent. This brings the next research idea:

VISION 6. *Have an FM-agent that could (i) involve human-in-the-loop and implement real-time monitoring to track model performance and detect shifts in user behavior or data distribution; (ii) use feedback from either human or AI to adjust the model (e.g. enhance model fairness when there are unexpected distribution shifts).*

Specifically, we may use a feedback loop to gather data on model predictions and user interactions, allowing for prompt detection of shifts in behavior [6, 84]. Once substantial shifts in user behaviors are detected, we could then involve human experts, potentially from partnering non-profit organizations, to review and provide feedback on model predictions and decisions. This feedback can then be used to guide model adjustments and improve its response to distribution changes. We may retrain or finetune the model using newly collected data, potentially weighting recent data more to better align with current trends and user behavior [8, 12, 40, 56, 95].

ACKNOWLEDGMENTS

The work is supported by Harvard HDSI funding.

REFERENCES

- [1] Joseph J Amon. 2020. Ending discrimination in healthcare. *Journal of the International AIDS Society* 23, 2 (2020).
- [2] Seyed Amir Hossein Aqajari, Ziyu Wang, Ali Tazarv, Sina Labbaf, Salar Jafar-lou, Brenda Nguyen, Nikil Dutt, Marco Levorato, and Amir M Rahmani. 2024. Enhancing Performance and User Engagement in Everyday Stress Monitoring: A Context-Aware Active Reinforcement Learning Approach. *arXiv preprint arXiv:2407.08215* (2024).
- [3] Lisa P Argyle, Ethan C Busby, Nancy Fulda, Joshua R Gubler, Christopher Rytting, and David Wingate. 2023. Out of one, many: Using language models to simulate human samples. *Political Analysis* 31, 3 (2023), 337–351.
- [4] Turgay Ayer, Can Zhang, Anthony Bonifonte, Anne C Spaulding, and Jagpreet Chhatwal. 2019. Prioritizing hepatitis C treatment in US prisons. *Operations Research* 67, 3 (2019), 853–873.
- [5] Rob Baltussen and Louis Niessen. 2006. Priority setting of health interventions: the need for multi-criteria decision analysis. *Cost effectiveness and resource allocation* 4 (2006), 1–9.
- [6] Nikhil Behari, Edwin Zhang, Yunfan Zhao, Aparna Taneja, Dheeraj Nagaraj, and Milind Tambe. 2024. A Decision-Language Model (DLM) for Dynamic Restless Multi-Armed Bandit Tasks in Public Health. *Neural Information Processing Systems* (2024).
- [7] Joshua Blumenstock. 2018. Don't forget people in the use of big data for development.
- [8] Rishi Bommasani, Drew A Hudson, Ehsan Adeli, Russ Altman, Simran Arora, Sydney von Arx, Michael S Bernstein, Jeannette Bohg, Antoine Bosselut, Emma Brunskill, et al. 2021. On the opportunities and risks of foundation models. *arXiv preprint arXiv:2108.07258* (2021).
- [9] Elizabeth Bondi, Lily Xu, Diana Acosta-Navas, and Jackson A Killian. 2021. Envisioning communities: a participatory approach towards AI for social good. In *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*. 425–436.
- [10] Weize Chen, Ziming You, Ran Li, Yitong Guan, Chen Qian, Chenyang Zhao, Cheng Yang, Ruobing Xie, Zhiyuan Liu, and Maosong Sun. 2024. Internet of agents: Weaving a web of heterogeneous agents for collaborative intelligence. *arXiv preprint arXiv:2407.07061* (2024).
- [11] Myra Cheng, Tiziano Piccardi, and Diyi Yang. 2023. CoMPoS: Characterizing and Evaluating Caricature in LLM Simulations. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*. 10853–10875.
- [12] Krzysztof Marcin Choromanski, Arijit Sehanobish, Han Lin, Yunfan Zhao, Eli Berger, Tetiana Parshakova, Alvin Pan, David Watkins, Tianyi Zhang, Valerii Likhoshesterov, et al. 2023. Efficient graph field integrators meet point clouds. In *International Conference on Machine Learning*. PMLR, 5978–6004.
- [13] Daniel D Corkill and Victor R Lesser. 1983. The Use of Meta-Level Control for Coordination in a Distributed Problem Solving Network. In *IJCAI*, Vol. 83. 748.
- [14] Josh Cows, Andreas Tsamados, Mariarosaria Taddeo, and Luciano Floridi. 2021. A definition, benchmark and database of AI for social good initiatives. *Nature Machine Intelligence* 3, 2 (2021), 111–115.
- [15] Maria De-Arteaga, William Herlands, Daniel B Neill, and Artur Dubrawski. 2018. Machine learning for the developing world. *ACM Transactions on Management Information Systems (TMIS)* 9, 2 (2018), 1–14.
- [16] Sarang Deo, Seyed Iravani, Tingting Jiang, Karen Smilowitz, and Stephen Samuelson. 2013. Improving health outcomes through better capacity allocation in a community-based chronic care model. *Operations Research* 61, 6 (2013), 1277–1294.
- [17] Heng Dong, Tonghan Wang, Jiayuan Liu, and Chongjie Zhang. 2022. Low-rank modular reinforcement learning via muscle synergy. *Advances in Neural Information Processing Systems* 35 (2022), 19861–19873.
- [18] Heng Dong, Junyu Zhang, Tonghan Wang, and Chongjie Zhang. 2023. Symmetry-aware robot design with structured subgroups. In *International Conference on Machine Learning*. PMLR, 8334–8355.
- [19] Jason Xiaotian Dou, Runxue Bao, and Wenxin Wei. 2022. Clinical Decision System using Machine Learning and Deep Learning: a Survey. (2022).
- [20] Jason Xiaotian Dou, Minxue Jia, Nika Zaslavsky, Mark Ebeid, Runxue Bao, Shiyi Zhang, Ke Ni, Paul Pu Liang, Haiyi Mao, and Zhi-Hong Mao. 2022. Learning more effective cell representations efficiently. In *NeurIPS 2022 Workshop on Learning Meaningful Representations of Life*.
- [21] Joshua K Dubrow. 2022. Local data and upstream reporting as sources of error in the administrative data undercount of Covid 19. *International Journal of Social Research Methodology* 25, 4 (2022), 471–476.
- [22] Cynthia Dwork, Aaron Roth, et al. 2014. The algorithmic foundations of differential privacy. *Foundations and Trends® in Theoretical Computer Science* 9, 3–4 (2014), 211–407.
- [23] Adam Elmachtoub, Vishal Gupta, and Yunfan Zhao. 2023. Balanced off-policy evaluation for personalized pricing. In *International Conference on Artificial Intelligence and Statistics*. PMLR, 10901–10917.
- [24] Adam N Elmachtoub, Henry Lam, Haofeng Zhang, and Yunfan Zhao. 2023. Estimate-then-optimize versus integrated-estimationoptimization: A stochastic dominance perspective. *arXiv preprint arXiv:2304.06833* (2023).
- [25] Luciano Floridi, Josh Cows, Monica Beltrametti, Raja Chatila, Patrice Chazerand, Virginia Dignum, Christoph Luetge, Robert Madelin, Ugo Pagallo, Francesca Rossi, et al. 2018. AI4People—an ethical framework for a good AI society: opportunities, risks, principles, and recommendations. *Minds and machines* 28 (2018), 689–707.
- [26] Francesca Foffano, Teresa Scantamburlo, and Atia Cortés. 2023. Investing in AI for social good: an analysis of European national strategies. *AI & society* 38, 2 (2023), 479–500.
- [27] Zohreh S Gatmiry, Ashkan Hafezalkotob, Roya Soltani, et al. 2021. Food web conservation vs. strategic threats: A security game approach. *Ecological Modelling* 442 (2021), 109426.
- [28] Michael R Genesereth and D Smith. 1983. An Overview of Meta-Level Architecture. In *AAAI*. 119–124.
- [29] Daniel Golovin, Andreas Krause, Beth Gardner, Sarah Converse, and Steve Morey. 2011. Dynamic resource allocation in conservation planning. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 25. 1331–1336.
- [30] Cleotilde Gonzalez, Javier F Lerch, and Christian Lebiere. 2003. Instance-based learning in dynamic decision making. *Cognitive Science* 27, 4 (2003), 591–635.
- [31] Lucia Gordon, Esther Rolf, and Milind Tambe. 2024. Combining Diverse Information for Coordinated Action: Stochastic Bandit Algorithms for Heterogeneous Agents. *arXiv preprint arXiv:2408.03405* (2024).
- [32] T Guo, X Chen, Y Wang, R Chang, S Pei, NV Chawla, O Wiest, and X Zhang. 2024. Large Language Model based Multi-Agents: A Survey of Progress and Challenges. In *33rd International Joint Conference on Artificial Intelligence (IJCAI 2024)*. IJCAI; Cornell arxiv.
- [33] Nan He, Hanyu Lai, Chenyang Zhao, Zirui Cheng, Juntao Pan, Ruoyu Qin, Ruofan Lu, Rui Lu, Yunchen Zhang, Gangming Zhao, et al. 2023. TeacherIm: Teaching to fish rather than giving the fish, language modeling likewise. *arXiv preprint arXiv:2310.19019* (2023).
- [34] Shengding Hu, Yuge Tu, Xu Han, Chaoqun He, Ganqu Cui, Xiang Long, Zhi Zheng, Yewei Fang, Yuxiang Huang, Weilin Zhao, et al. 2024. Minicpm: Unveiling the potential of small language models with scalable training strategies. *arXiv preprint arXiv:2404.06395* (2024).
- [35] Yuelu Ji, Zhuochun Li, Rui Meng, and Daqing He. 2024. ReasoningRank: Teaching Student Models to Rank through Reasoning-Based Knowledge Distillation. *arXiv preprint arXiv:2410.05168* (2024).
- [36] Yuelu Ji, Wenhe Ma, Sonish Sivarajkumar, Hang Zhang, Eugene Mathew Sadhu, Zhuochun Li, Xizhi Wu, Shyam Visweswaran, and Yanshan Wang. 2024. Mitigating the Risk of Health Inequity Exacerbated by Large Language Models. *arXiv preprint arXiv:2410.05180* (2024).
- [37] Yuelu Ji, Zeshui Yu, and Yanshan Wang. 2024. Assertion Detection in Clinical Natural Language Processing Using Large Language Models. In *2024 IEEE 12th International Conference on Healthcare Informatics (ICHI)*. 242–247. <https://doi.org/10.1109/ICHI61247.2024.00039>
- [38] Yipeng Kang, Tonghan Wang, Qianlan Yang, Xiaoran Wu, and Chongjie Zhang. 2022. Non-linear coordination graphs. *Advances in Neural Information Processing Systems* 35 (2022), 25655–25666.
- [39] Jacob Devlin Ming-Wei Chang Kenton and Lee Kristina Toutanova. 2019. Bert: Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of naacl-HLT*, Vol. 1. 2.
- [40] Linghai Kong, Haorui Wang, Wenhao Mu, Yuanqi Du, Yuchen Zhuang, Yifei Zhou, Yue Song, Rongzhi Zhang, Kai Wang, and Chao Zhang. 2024. Aligning Large Language Models with Representation Editing: A Control Perspective. *Advances in Neural Information Processing Systems* (2024).
- [41] Margaret E Kruk, Anna D Gage, Catherine Arsenault, Keely Jordan, Hannah H Leslie, Sanam Roder-DeWan, Olusoji Adeyi, Pierre Barker, Bernadette Daelmans, Svetlana V Doubova, et al. 2018. High-quality health systems in the Sustainable Development Goals era: time for a revolution. *The Lancet global health* 6, 11 (2018), e1196–e1252.
- [42] Roberta Kwok et al. 2019. AI empowers conservation biology. *Nature* 567, 7746 (2019), 133–134.
- [43] Shuoqi Li, Han Xu, and Haipeng Chen. 2024. Focused ReAct: Improving ReAct through Reiterate and Early Stop. *arXiv preprint arXiv:2410.10779* (2024).
- [44] Yinghao Li, Linghai Kong, Yuanqi Du, Yue Yu, Yuchen Zhuang, Wenhao Mu, and Chao Zhang. 2023. Muben: Benchmarking the uncertainty of molecular representation models. *Transactions on Machine Learning Research* (2023).
- [45] Yuchen Li, Siyu Teng, Lingxi Li, Zhe Xuanyuan, and Long Chen. 2023. Foundation Models for Mining 5.0: Challenges, Frameworks, and Opportunities. In *2023 IEEE 3rd International Conference on Digital Twins and Parallel Intelligence (DTPi)*. 1–6. <https://doi.org/10.1109/DTPi59677.2023.10365454>
- [46] Haiyi Mao, Minxue Jia, Jason Xiaotian Dou, Haotian Zhang, and Panayiotis V Benos. 2022. COEM: cross-modal embedding for metacell identification. *arXiv preprint arXiv:2207.07734* (2022).
- [47] Haiyi Mao, Hongfu Liu, Jason Xiaotian Dou, and Panayiotis V Benos. 2022. Towards cross-modal causal structure and representation learning. In *Machine Learning for Health*. PMLR, 120–140.
- [48] Julia M Markel, Steven G Opferman, James A Landay, and Chris Piech. 2023. Gpteach: Interactive training with gpt-based students. In *Proceedings of the tenth acm conference on learning@scale*. 226–236.

- [49] Aditya Mate, Lovish Madaan, Aparna Taneja, Neha Madhiwalla, Shresth Verma, Gargi Singh, Aparna Hegde, Pradeep Varakantham, and Milind Tambe. 2022. Field study in deploying restless multi-armed bandits: Assisting non-profits in improving maternal and child health. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 36. 12017–12025.
- [50] Tasquia Mizan and Sharareh Taghipour. 2022. Medical resource allocation planning by integrating machine learning and optimization models. *Artificial Intelligence in Medicine* 134 (2022), 102430.
- [51] Michael Moor, Oishi Banerjee, Zahra Shakeri Hossein Abad, Harlan M Krumholz, Jure Leskovec, Eric J Topol, and Pranav Rajpurkar. 2023. Foundation models for generalist medical artificial intelligence. *Nature* 616, 7956 (2023), 259–265.
- [52] Siddharth Nishtala, Lovish Madaan, Aditya Mate, Harshvardhan Kamarthi, Anirudh Grama, Divy Thakkar, Dhyanesh Narayanan, Suresh Chaudhary, Neha Madhiwalla, Ramesh Padmanabhan, et al. 2021. Selective intervention planning using restless multi-armed bandits to improve maternal and child health outcomes. *arXiv preprint arXiv:2103.09052* (2021).
- [53] Joon Sung Park, Joseph O’Brien, Carrie Jun Cai, Meredith Ringel Morris, Percy Liang, and Michael S Bernstein. 2023. Generative agents: Interactive simulacra of human behavior. In *Proceedings of the 36th annual ACM symposium on user interface software and technology*. 1–22.
- [54] Sarita Paudel, Benjamin E Warner, Renwei Wang, Jennifer Adams-Haduch, Alex S Reznik, Jason Dou, Yufei Huang, Yu-Tang Gao, Woon-Puay Koh, Alan Bäckholm, et al. 2022. Serologic profiling using an Epstein-Barr virus mammalian expression library identifies EBNA1 IgA as a prediagnostic marker for nasopharyngeal carcinoma. *Clinical Cancer Research* 28, 23 (2022), 5221–5230.
- [55] David Paulus, Gerdien de Vries, Marijn Janssen, and Bartel Van de Walle. 2023. Reinforcing data bias in crisis information management: The case of the Yemen humanitarian response. *International Journal of Information Management* 72 (2023), 102663.
- [56] Dazhi Peng and Hangrui Cao. 2024. E-tamba: Efficient transformer-mamba layer transplantation. In *NeurIPS 2024 Workshop on Fine-Tuning in Modern Machine Learning: Principles and Scalability*.
- [57] Andrew Perrault, Fei Fang, Arunesh Sinha, and Milind Tambe. 2020. Artificial Intelligence for Social Impact: Learning and Planning in the Data-to-Deployment Pipeline. *AI Mag.* 41, 4 (2020), 3–16. <https://doi.org/10.1609/AIMAG.V41i4.5296>
- [58] Yuji Roh, Geon Heo, and Steven Euijong Whang. 2019. A survey on data collection for machine learning: a big data-ai integration perspective. *IEEE Transactions on Knowledge and Data Engineering* 33, 4 (2019), 1328–1347.
- [59] Paul S Rosenbloom, John E Laird, and Allen Newell. 1988. Meta-levels in Soar. In *Meta-Level Architectures and Reflection*. Elsevier Science Publishers BV Amsterdam, 227–240.
- [60] Nithya Sambasivan and Rajesh Veeraraghavan. 2022. The deskilling of domain expertise in AI development. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems*. 1–14.
- [61] Shibani Santurkar, Esin Durmus, Faisal Ladhak, Cinoo Lee, Percy Liang, and Tatsunori Hashimoto. 2023. Whose opinions do language models reflect?. In *International Conference on Machine Learning*. PMLR, 29971–30004.
- [62] Roderick Seow, Yunfan Zhao, Duncan Wood, Milind Tambe, and Cleotilde Gonzalez. 2024. Improving the Prediction of Individual Engagement in Recommendations Using Cognitive Models. *Workshop on Health Recommender Systems co-located with ACM RecSys 2024* (2024).
- [63] Sonia Jawaid Shaikh. 2020. Artificial Intelligence and Resource Allocation in Healthcare: The Process-Outcome Divide in Perspectives on Moral Decision-Making. In *AI4SG@ AAAI Fall Symposium*.
- [64] Zheyuan Ryan Shi, Claire Wang, and Fei Fang. 2020. Artificial intelligence for social good: A survey. *arXiv preprint arXiv:2001.01818* (2020).
- [65] Eric Shieh, Bo An, Rong Yang, Milind Tambe, Craig Baldwin, Joseph DiRenzo, Ben Maule, and Garrett Meyer. 2012. Protect: A deployed game theoretic system to protect the ports of the united states. In *Proceedings of the 11th international conference on autonomous agents and multiagent systems-volume 1*. 13–20.
- [66] Arunesh Sinha, Fei Fang, Bo An, Christopher Kiekintveld, and Milind Tambe. 2018. Stackelberg security games: Looking beyond a decade of success. *IJCAI*.
- [67] Haotian Sun, Yuchen Zhuang, Linghai Kong, Bo Dai, and Chao Zhang. 2023. Adaplanner: Adaptive planning from feedback with language models. *Advances in Neural Information Processing Systems* (2023).
- [68] Nenad Tomašev, Julien Cornéise, Frank Hutter, Shakir Mohamed, Angela Picciariello, Bec Connelly, Danielle CM Belgrave, Daphne Ezer, Fanny Cachat van der Haert, Frank Mugisha, et al. 2020. AI for social good: unlocking the opportunity for positive impact. *Nature Communications* 11, 1 (2020), 2468.
- [69] Shresth Verma, Niclas Boehmer, Linghai Kong, and Milind Tambe. 2024. Balancing Act: Prioritization Strategies for LLM-Designed Restless Bandit Rewards. (2024). *arXiv:2408.12112*
- [70] Shresth Verma, Yunfan Zhao, Sanket Shah, Niclas Boehmer, Aparna Taneja, and Milind Tambe. [n.d.]. Group Fairness in Predict-Then-Optimize Settings for Restless Bandits. In *The 40th Conference on Uncertainty in Artificial Intelligence*.
- [71] Vijay Viswanathan, Chenyang Zhao, Amanda Bertsch, Tongshuang Wu, and Graham Neubig. 2023. Prompt2Model: Generating Deployable Models from Natural Language Instructions. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*. 413–421.
- [72] Fei Wang and Anita Preininger. 2019. AI in health: state of the art, challenges, and future directions. *Yearbook of medical informatics* 28, 01 (2019), 016–026.
- [73] Haorui Wang, Marta Skreta, Cher-Tian Ser, Wenhao Gao, Linghai Kong, Felix Streith-Kalthoff, Chenru Duan, Yuchen Zhuang, Yue Yu, Yanqiao Zhu, et al. 2024. Efficient Evolutionary Search over Chemical Space with Large Language Models. *arXiv preprint arXiv:2406.16976* (2024).
- [74] Jiaqi Wang, Zihao Wu, Yiwei Li, Hanqi Jiang, Peng Shu, Enze Shi, Huawen Hu, Chong Ma, Yiheng Liu, Xuhui Wang, et al. 2024. Large language models for robotics: Opportunities, challenges, and perspectives. *arXiv preprint arXiv:2401.04334* (2024).
- [75] Kai Wang, Shresth Verma, Aditya Mate, Sanket Shah, Aparna Taneja, Neha Madhiwalla, Aparna Hegde, and Milind Tambe. 2023. Scalable decision-focused learning in restless multi-armed bandits with application to maternal and child health. In *Proceedings of the AAAI Conference on Artificial Intelligence*, Vol. 37. 12138–12146.
- [76] Tonghan Wang, Heng Dong, Victor Lesser, and Chongjie Zhang. 2020. ROMA: Multi-Agent Reinforcement Learning with Emergent Roles. In *Proceedings of the 37th International Conference on Machine Learning*.
- [77] Tonghan Wang, Tarun Gupta, Anuj Mahajan, Bei Peng, Shimon Whiteson, and Chongjie Zhang. 2020. RODE: Learning Roles to Decompose Multi-Agent Tasks. In *International Conference on Learning Representations*.
- [78] Yanan Wang, Yong Ge, Li Li, Rui Chen, and Tong Xu. 2020. M3rec: An off-line meta-level model-based reinforcement learning approach for cold-start recommendation. In *Proceedings of the 33th International Conference on Neural Information Processing Systems, NIPS, Vol. 20*.
- [79] Ziyu Wang, Anil Kanduri, Seyed Amir Hossein Aqajari, Salar Jafarlou, Sanaz R Mousavi, Pasi Liljeberg, Shaista Malik, and Amir M Rahmani. 2024. Ecg unveiled: Analysis of client re-identification risks in real-world ecg datasets. In *2024 IEEE 20th International Conference on Body Sensor Networks (BSN)*. IEEE, 1–4.
- [80] Ziyu Wang, Hao Li, Di Huang, and Amir M Rahmani. 2024. HealthQ: Unveiling Questioning Capabilities of LLM Chains in Healthcare Conversations. *arXiv preprint arXiv:2409.19487* (2024).
- [81] Ziyu Wang, Nanqing Luo, and Pan Zhou. 2020. GuardHealth: Blockchain empowered secure data management and Graph Convolutional Network enabled anomaly detection in smart healthcare. *J. Parallel and Distrib. Comput.* 142 (2020), 1–12.
- [82] Ziyu Wang, Zhongqi Yang, Iman Azimi, and Amir M. Rahmani. 2024. Differential Private Federated Transfer Learning for Mental Health Monitoring in Everyday Settings: A Case Study on Stress Detection. In *Proceedings of the 46th Annual International Conference of the IEEE Engineering in Medicine and Biology Society (EMBC)* (USA).
- [83] Chaojun Xiao, Zhengyan Zhang, Chenyang Song, Dazhi Jiang, Feng Yao, Xu Han, Xiaozhi Wang, Shuo Wang, Yufei Huang, Guanyu Lin, et al. 2024. Configurable Foundation Models: Building LLMs from a Modular Perspective. *arXiv preprint arXiv:2409.02877* (2024).
- [84] Guojun Xiong, Ujwal Dinesha, Debajoy Mukherjee, Jian Li, and Srinivas Shakkottai. 2024. DOPL: Direct Online Preference Learning for Restless Bandits with Preference Feedback. *arXiv:2410.05527 [cs.LG]* <https://arxiv.org/abs/2410.05527>
- [85] Guojun Xiong and Jian Li. 2024. Provably Efficient Reinforcement Learning for Adversarial Restless Multi-Armed Bandits with Unknown Transitions and Bandit Feedback. *arXiv preprint arXiv:2405.00950* (2024).
- [86] Guojun Xiong, Shufan Wang, Jian Li, and Rahul Singh. 2024. Whittle Index-Based Q-Learning for Wireless Edge Caching With Linear Function Approximation. *IEEE/ACM Transactions on Networking* (2024).
- [87] Guojun Xiong, Shufan Wang, Gang Yan, and Jian Li. 2023. Reinforcement learning for dynamic dimensioning of cloud caches: A restless bandit approach. *IEEE/ACM Transactions on Networking* 31, 5 (2023), 2147–2161.
- [88] Haifeng Xu, Long Tran-Thanh, and Nick Jennings. 2016. Playing repeated security games with no prior knowledge. In *AAMAS’16: Proceedings of the 2016 International Conference on Autonomous Agents & Multiagent Systems*. ACM Press, 104–112.
- [89] Han Xu, Xingyuan Wang, and Haipeng Chen. 2024. Towards Real-Time and Personalized Code Generation. In *Proceedings of the 33rd ACM International Conference on Information and Knowledge Management*. 5568–5569.
- [90] Han Xu, Jingyang Ye, Yutong Li, and Haipeng Chen. 2024. Can Speculative Sampling Accelerate ReAct Without Compromising Reasoning Quality?. In *The Second Tiny Papers Track at ICLR 2024*. <https://openreview.net/forum?id=42b9hJrlpX>
- [91] Han Xu, Ruining Zhao, Jindong Wang, and Haipeng Chen. 2024. RESTful-Llama: Connecting User Queries to RESTful APIs. In *Proceedings of the 2024 Conference on Empirical Methods in Natural Language Processing: Industry Track*. 1433–1443.
- [92] Liming Xu, Sara Almahri, Stephen Mak, and Alexandra Brintrup. 2023. Multi-Agent Systems and Foundation Models Enable Autonomous Supply Chains: Opportunities and Challenges. *Available at SSRN 4695075* (2023).
- [93] Lily Xu, Arpita Biswas, Fei Fang, and Milind Tambe. 2022. Ranked prioritization of groups in combinatorial bandit allocation. *arXiv preprint arXiv:2205.05659* (2022).

- [94] Meike Zehlike, Francesco Bonchi, Carlos Castillo, Sara Hajian, Mohamed Megahed, and Ricardo Baeza-Yates. 2017. Fa* ir: A fair top-k ranking algorithm. In *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*. 1569–1578.
- [95] Chenyang Zhao, Xueying Jia, Vijay Viswanathan, Graham Neubig, and Tongshuang Wu. [n.d.]. Self-Guide: Better Task-Specific Instruction Following via Self-Synthetic Finetuning. In *First Conference on Language Modeling*.
- [96] Yunfan Zhao, Nikhil Behari, Edward Hughes, Edwin Zhang, Dheeraj Nagaraj, Karl Tuyls, Aparna Taneja, and Milind Tambe. 2024. Towards a Pretrained Model for Restless Bandits via Multi-arm Generalization. *IJCAI*.
- [97] Yunfan Zhao, Nikhil Behari, Edward Hughes, Edwin Zhang, Dheeraj Nagaraj, Karl Tuyls, Aparna Taneja, and Milind Tambe. 2024. Towards Zero Shot Learning in Restless Multi-armed Bandits. In *Proceedings of the 23rd International Conference on Autonomous Agents and Multiagent Systems*. 2618–2620.
- [98] Yunfan Zhao, Alvin Pan, Krzysztof Marcin Choromanski, Deepali Jain, and Vikas Sindhwani. [n.d.]. Implicit Two-Tower Policies. In *5th Workshop on practical ML for limited/low resource settings*.
- [99] Yunfan Zhao, Tonghan Wang, Dheeraj Nagaraj, Aparna Taneja, and Milind Tambe. 2024. The Bandit Whisperer: Communication Learning for Restless Bandits. *arXiv preprint arXiv:2408.05686* (2024).