
Building Socially-Equitable Public Models

Yejia Liu¹ Jianyi Yang¹ Pengfei Li¹ Tongxin Li² Shaolei Ren¹

Abstract

Public models offer predictions to a variety of downstream tasks and have played a crucial role in various AI applications, showcasing their proficiency in accurate predictions. However, the exclusive emphasis on prediction accuracy may not align with the diverse end objectives of downstream agents. Recognizing the public model’s predictions as a service, we advocate for integrating the objectives of downstream agents into the optimization process. Concretely, to address performance disparities and foster fairness among heterogeneous agents in training, we propose a novel Equitable Objective. This objective, coupled with a policy gradient algorithm, is crafted to train the public model to produce a more equitable/uniform performance distribution across downstream agents, each with their unique concerns. Both theoretical analysis and empirical case studies have proven the effectiveness of our method in advancing performance equity across diverse downstream agents utilizing the public model for their decision-making. Codes and datasets are released at <https://github.com/Ren-Research/Socially-Equitable-Public-Models>.

1. Introduction

Public models whose outputs are utilized by multiple agents have become essential building blocks for multiple AI applications such as climate modeling and traffic prediction. These models undergo training on extensive datasets and are tailored for specific domains, making them highly effective in generating accurate predictions (Nguyen et al., 2023; Bommasani et al., 2021; Shah et al., 2022). Their accessibility and availability to the public enable the widespread

utilization by diverse downstream agents for individual business goals (Yang et al., 2023). However, it is important to note that exclusive reliance on prediction accuracy may not be ideal when serving a diverse range of downstream agents, each with unique business objectives. Consider a scenario where a public model predicts disease outbreaks across different regions. While accuracy is pivotal, optimizing the allocation of healthcare resources—ensuring sufficient medical supplies, personnel, and preventive measures—takes precedence, based on the general prediction provided by the public model.

We therefore suggest taking into account the impact of a public model’s prediction on downstream agents, rather than solely focusing on minimizing prediction errors during training. A closely related topic is *decision-focused learning*, which involves incorporating domain-specific constraints and/or objectives into the learning algorithm (Johnson-Yu et al., 2023; Wilder et al., 2020). However, the majority of existing decision-focused learning works only address a single task or agent, rendering them barely applicable to the challenge faced by public models, which deal with diverse downstream agents with their varied decision-making objectives (Mandi et al., 2023). Additionally, in the current decision-focused learning framework, performance disparities can arise, with certain agents consistently facing inferior outcomes. For instance, this may happen when some agents have limited training data availability, while others have access to abundant and various datasets.

We view the prediction provided by a public model as a service for diverse downstream agents. As a service provider, prioritizing accuracy is crucial, but ensuring high-quality service for all users, given their diverse concerns, is even more vital. Unfairly benefiting or disadvantaging model performance on specific agents is unjust. While machine learning fairness studies primarily concentrate on achieving accuracy balance among protected groups with sensitive characteristics (Barocas et al., 2023; Pessach & Shmueli, 2022), we introduce a different fairness perspective centered on ensuring performance equity/uniformity across downstream agents with different decision processes. Recent works have proposed a related concept referred to as the “good-intent” fairness, primarily focused on preventing overfitting to any specific device in federated learning (Mohri et al., 2019; Li et al., 2020). However, its scope is limited

¹University of California, Riverside, United States ²The Chinese University of Hong Kong, Shenzhen, China. Correspondence to: Shaolei Ren <sren@ece.ucr.edu>, Yejia Liu <yliu807@ucr.edu>.

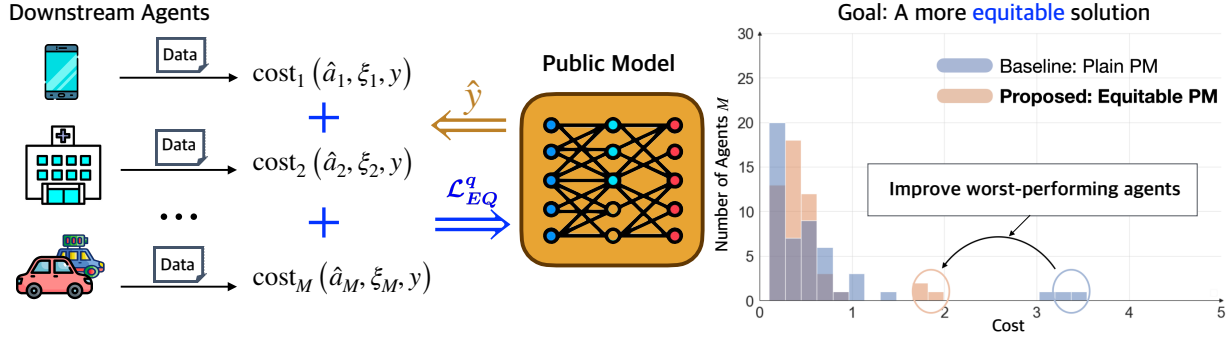


Figure 1: The EQUITABLE PM leads to a fairer solution by fostering a more equitable/uniform performance distribution across downstream agents. The embedded Equitable Objective $\mathcal{L}_{EQ}^q$ directly accounts for the decision costs across diverse downstream agents that use the prediction $\hat{y}$ from the public model f for making informed decisions via agent-specific decision processes and actions $(\hat{a}_1, \dots, \hat{a}_M)$.

to maximizing the performance of the worst-performing devices without accounting for the decision processes and objectives of diverse downstream agents.

In this work, we propose the Equitable Objective, inspired by the α -fairness in resource allocation (Altman et al., 2008), to tackle fairness concerns while considering decision-makings of downstream agents in the development of a public model. The objective minimizes an aggregated reweighted loss, parameterized by q , prioritizing the optimization of worse costs—assigning higher relative weights to agents with higher downstream costs when leveraging predictions from a public model. As shown in the motivating example illustrated in Figure 1, the proposed approach leads to a more equitable performance among heterogeneous agents compared to the baseline, which solely minimizes prediction errors through MSE loss.

Contributions. We consider a novel setting and propose an Equitable Objective to ensure performance equity/uniformity across downstream agents leveraging a public model for decision-making. We then present an algorithm to optimize the proposed Equitable Objective, which is applicable to both differentiable and non-differentiable downstream cost functions. Additionally, we provide theoretical results guaranteeing performance equity/uniformity of the proposed approach, along with insights into generalization bounds. Empirically, we demonstrate through case studies using real-world datasets that our approach leads to a more equitable/uniform cost distribution among downstream agents under various settings.

2. Problem Formulation

Consider a public model, denoted by $f : \mathbf{X} \times \Theta \rightarrow \mathbf{Y}$ where $\mathbf{X}$ is an input space, Θ is a set of parameters, and $\mathbf{Y}$ is an output space. The inputs $x \in \mathbf{X}$ are features shared by multiple downstream tasks. For any $x \in \mathbf{X}$ and $\theta \in$

Θ , we write $\hat{y} := f(x; \theta)$ as a prediction from the public model f . A significant emphasis in prevalent public model training is on minimizing prediction errors and achieving high accuracy (Bommasani et al., 2021). However, the loss function used for model training can be easily misaligned with the ultimate goal, which is to optimize decision-making when utilized by diverse downstream agents. We therefore suggest *incorporating downstream agents' costs* into the objective formulation.

Suppose that there are M heterogeneous downstream agents employing the public model f for decision-making in a stochastic environment. Each agent m possesses a context variable ξ_m , which can either represent public shared features like local weather conditions or encapsulate unique features of downstream agents. By following a policy π_m , each agent generates an action w.r.t. the input, denoted as $\hat{a}_m(\theta) := \pi_m(\hat{y}, \xi_m)$. The resulting action $\hat{a}_m(\theta)$ taken by the agent m would incur a cost, represented as $\text{cost}_m(\hat{a}_m(\theta), \xi_m, y)$.

To address the decision cost of downstream agents, a straightforward approach is to formulate the objective as

$$\min_{\theta} \sum_{m=1}^M \mathbb{E} [\text{cost}_m(\hat{a}_m(\theta), \xi_m, y) - \text{cost}_m(a_m, \xi_m, y)],$$

where $a_m = \arg \min_{a \in \mathbf{A}} \text{cost}_m(a, \xi_m, y)$ and $\mathbf{A}$ represents the action space. That is, the objective is to minimize the total expected *regret* (i.e., the cost of decisions made based on predicted $\hat{y}$ minus the cost of decisions based on the true y) for all the M downstream agents due to the public model's potential prediction errors.¹

In an illustrative example where the public model f optimizes traffic signal timings, the standard accuracy goal is to

¹Our study can be easily generalized to incorporate an additional weight for the expected regret of each agent.

minimize delays, $\min_{\theta} E[(y - \hat{y})^2]$, where y is actual traffic conditions and $\hat{y}$ is predicted traffic flow. In reality, the transportation system involves diverse downstream stakeholders with unique concerns: commuters prioritize travel time and fuel consumption, public services focus on schedule, and environmental regulators are concerned with carbon emissions. Each party faces decision costs from the actions it takes based on the model's predictions $\hat{y}$. Thus, using an objective encompassing diverse costs from downstream agents can explicitly incorporate their concerns.

Nevertheless, due to the heterogeneity of agents such as various data biases, solely minimizing the total cost objective can result in significant performance disparities among downstream agents. Consequently, certain agents may consistently experience the poorest performance when using the prediction provided by the public model compared to other agents. For example, the trained model may exhibit a preference towards the agents with greater numbers of data samples. This inequity in performance highlights concerns regarding the fairness of the services these agents receive when viewing the prediction from the public model as a shared "resource" serving diverse downstream agents.

3. Fair Public Model for Downstream Agents

To achieve fairness for the downstream agents with different decision processes, we propose EQUITABLE PM, which seeks to optimize a novel Equitable Objective.

3.1. Defining Fairness: An Equitable Objective

We now introduce the Equitable Objective to address fairness concerns in the context of diverse downstream costs across different agents. By drawing inspiration from the α -fairness resource allocation (Altman et al., 2008; Jang & Yang, 2022; Li et al., 2020), we propose an objective to promote performance equity/uniformity parameterized by $q \geq 0$. Both theoretical proofs (Section 4) and empirical case studies (Section 5) have shown that the use of the Equitable Objective leads to a more equitable/uniform performance distribution across downstream agents.

The Equitable Objective aims to minimize the aggregated cost incurred by downstream agents, parameterized by q , when utilizing the prediction from the public model f , as shown in Eq. (1),

$$\min_{\theta} \mathcal{J}_{EQ}^q(\theta) := \sum_{m=1}^M \mathbb{E}^{q+1} [\text{cost}_m(\hat{a}_m(\theta), \xi_m, y) - \text{cost}_m(a_m, \xi_m, y)], \quad (1)$$

where the hyperparameter $q \geq 0$ promotes performance equity among different agents. Specifically, when q is set larger, the minimization process will take into account the agent of worst performance to a greater extent.

To train a public model, we need to empirically approximate (1) with training data samples. Let $\mathcal{D}_m = \{x_{m,i}, y_{m,i}, \xi_{m,i} | i \in [N_m]\}$ be the dataset of the agent m , where N_m is the number of data examples in the agent m . Note that the public model's input features may still vary among different agents (e.g., a public carbon-intensity prediction model uses location-specific features to predict the local grid's carbon intensity). Thus, for two different agents m_1 and m_2 , the public variables $\{x_{m_1,i}, y_{m_1,i}\}$ and $\{x_{m_2,i}, y_{m_2,i}\}$ can be identical or different depending on factors such as whether they are collected at the same time and/or location. Given the datasets $\mathcal{D}_1, \dots, \mathcal{D}_M$, we can approximate the expectation $\mathcal{J}_{EQ}^q(\theta)$ in Eq. (1) with the empirical loss $\mathcal{L}_{EQ}^q(\theta)$ defined in Eq. (2),

$$\min_{\theta} \mathcal{L}_{EQ}^q(\theta) := \sum_{m=1}^M \left[\left(\frac{1}{N_m} \sum_{i=1}^{N_m} C_{m,i} \right)^{q+1} \right], \quad (2)$$

where we denote $C_{m,i} = \text{cost}_m(\hat{a}_{m,i}(\theta), \xi_{m,i}, y_{m,i}) - \text{cost}_m(a_{m,i}, \xi_{m,i}, y_{m,i})$ as the regret regarding the i th sample of agent m .

In practice, the public model developer may not always have direct access to the costs of all the downstream agents for training. In such cases, it can generate synthetic downstream agents by modeling their decision processes based on, e.g., utility maximization or cost minimization, for the target application. Additionally, annotation-sample efficient methods like task programming (Sun et al., 2021) can also help model the downstream decision processes.

It is worth noting that we have also proposed a more general objective in Appendix A.5, which combines $\mathcal{L}_{EQ}^q(\theta)$ with the public model's prediction loss $\mathcal{L}_f$ via a balancing hyperparameter $\beta \in [0, 1]$, allowing for a more nuanced control over the fairness-accuracy trade-off in optimization. In the subsequent text, we denote $\mathcal{L}_{EQ}^q(\theta)$ as $\mathcal{L}^q(\theta)$ for simplicity.

Our proposed $\mathcal{L}^q(\theta)$ ensures that the public model's predictions consider the diverse concerns of downstream agents. The trade-offs introduced by adjusting q contribute to a fairer distribution of performance across agents, fostering an equitable decision-making environment. In the subsequent sections, we provide details and algorithms to train a public model with the Equitable Objective.

3.2. Training Public Model: EQUITABLE PM

The difficulties in training a public model vary depending on the cost functions. When the cost functions are differentiable, it is feasible to calculate the gradient based on the chain rule. By back propagation, we can get the gradient as

$$\nabla_{\theta} \mathcal{L}^q(\theta) = \sum_{m=1}^M \sum_{i=1}^{N_m} \nabla_{C_{m,i}} \mathcal{L}^q \nabla_{\hat{a}_{m,i}} C_{m,i} \nabla_{\hat{y}_{m,i}} \hat{a}_{m,i} \nabla_{\theta} \hat{y}_{m,i},$$

Algorithm 1 EQUITABLE PM

Input: Training dataset, learning rate α
 Initialize the parameters θ
for each batch $k \in [K]$ **do**
 Obtain $\hat{y}_{m,k,i}$ by the public model $f(\cdot; \theta)$
 Compute the cost regret $C_{m,k,i}$ for the example $(x_{m,k,i}, y_{m,k,i}, \xi_{m,k,i})$ in batch k for $m \in [1, \dots, M]$
 Compute the gradient $\nabla_{\theta} \mathcal{L}_k^q(\theta)$ for batch k by Eq. (4).
 Update the parameter $\theta \leftarrow \theta - \alpha \nabla_{\theta} \mathcal{L}_k^q(\theta)$
end for

where we denote the regret of the i th sample of agent m as $C_{m,i} = \text{cost}_m(\hat{a}_{m,i}, \xi_{m,i}, y_{m,i}) - \text{cost}_m(a_{m,i}, \xi_{m,i}, y_{m,i})$.

Nonetheless, the training becomes significantly more challenging when the cost function is non-differentiable w.r.t. the $\hat{y}$. The non-differentiable cost function is prevalent for many practical downstream tasks. For example, some downstream tasks are combinatorial optimization problems with discrete actions (Wilder et al., 2019). Thus, a training method that does not rely on differentiable cost functions is critically needed for public models.

One possible method is to learn a differentiable model to approximate the cost function by observing the evolution of the sequence of actions and costs (Moerland et al., 2023; Yu et al., 2020). However, this method suffers from potentially inaccurate modeling of dynamic environments (Agarwal et al., 2023; Malik et al., 2019). Therefore, we can adopt a model-free approach, such as black-box optimization, which requires fewer assumptions about the underlying system (Agarwal et al., 2023). In our context, we choose the policy gradient (PG) algorithm, falling into the category of model-free approaches, in favor of its natural exploration-exploitation trade-off (Bhandari & Russo, 2024; Peters & Schaal, 2006). We next present the process of using PG to optimize $\mathcal{L}^q(\theta)$.

PG for training a public model differs notably from the standard PG algorithm. Due to the non-separable Equitable Objective in Eq. (2), the supervision loss hinges on the average regret of each agent m . Thus, we use a batch-based training approach. At each training step, we employ a probabilistic public model $\sigma_{\theta}(\hat{y} | x)$ to sample $\hat{y}_{m,i}$ given inputs $x_{m,i}, i \in [1, \dots, B_m]$ from a batch of B_m samples. By this way, we obtain an equitable loss expressed as $\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1}$. To utilize the equitable batch loss for supervising the training of the public model, we reformulate the original objective as

$$\mathbb{E}[\mathcal{L}^q(\theta)] = \mathbb{E}_{(X,Y,\hat{Y},\Xi) \sim p_{\theta}} \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right],$$

where p_{θ} is the joint distribution of the random variables

$X = [x_{m,i} \mid m \in [1, \dots, M], i \in [1, \dots, B_m]]$, $Y = [y_{m,i} \mid m \in [1, \dots, M], i \in [1, \dots, B_m]]$, $\hat{Y} = [\hat{y}_{m,i} \mid m \in [1, \dots, M], i \in [1, \dots, B_m]]$, and $\Xi = [\xi_{m,i} \mid m \in [1, \dots, M], i \in [1, \dots, B_m]]$, which relies on the probabilistic public model $\sigma_{\theta}(\hat{y} | x)$. The gradient of $\mathbb{E}[\mathcal{L}^q(\theta)]$ with respect to θ is given by

$$\nabla_{\theta} \mathbb{E}[\mathcal{L}^q(\theta)] = \mathbb{E}_{(X,Y,\hat{Y},\Xi) \sim p_{\theta}} \left\{ \left[\sum_{m=1}^M \sum_{i=1}^{B_m} \nabla_{\theta} \log \sigma_{\theta}(\hat{y}_{m,i} | x_{m,i}) \right] \cdot \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] \right\}, \quad (3)$$

whose detailed derivation can be found in Appendix A.1.

Given a training dataset with K batches, we can get an empirical approximation of the expected gradient in Eq. (3) as follows

$$\nabla_{\theta} \mathcal{L}^q(\theta) = \frac{1}{K} \sum_{k=1}^K \nabla_{\theta} \mathcal{L}_k^q(\theta), \quad (4)$$

$$\text{where } \nabla_{\theta} \mathcal{L}_k^q(\theta) = \left[\sum_{m=1}^M \sum_{i=1}^{B_m} \nabla_{\theta} \log \sigma_{\theta}(\hat{y}_{m,k,i} | x_{m,k,i}) \right] \cdot \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,k,i} \right)^{q+1} \right].$$

In summary, the training steps using PG to minimize $\mathcal{L}_{\theta}^q$ can be outlined as in the Algorithm 1. During inference, the public model is its deterministic counterpart expressed as $f(x; \theta) := \arg \max_{\hat{y}} \sigma_{\theta}(\hat{y} | x)$. In subsequent texts, we refer to our proposed method as the EQUITABLE PM.

4. Theoretical Analysis

4.1. Performance Equity/Uniformity

In this section, we provide the theoretical justification that the proposed Equitable Objective can promote greater equity/uniformity in the performance distribution across downstream tasks with proofs in Appendix A.2. We use $C_m = \frac{1}{N_m} \sum_{i=1}^{N_m} C_{m,i}$ to denote the performance of the m -th downstream agent. We here adopt *variance* and *entropy* to measure the uniformity of the performance distribution across downstream tasks.

Definition 4.1. (*Equity by Variance*) The performance distribution of M downstream agents $\{C_1(\theta), \dots, C_M(\theta)\}$ is more equitable/uniform under solution θ than θ' if

$$\text{Var}(C_1(\theta), \dots, C_M(\theta)) < \text{Var}(C_1(\theta'), \dots, C_M(\theta')), \quad (5)$$

where Var represents the *variance* of performance.

Definition 4.2. (*Equity by Entropy*) The performance distribution of M downstream agents $\{C_1(\theta), \dots, C_M(\theta)\}$ is

more equitable/uniform under solution θ than θ' if the entropy of the normalized performance distribution satisfies

$$\mathbb{H}_{\text{norm}}(C(\theta)) \geq \mathbb{H}_{\text{norm}}(C(\theta')), \quad (6)$$

where $\mathbb{H}_{\text{norm}}(C(\theta))$ is expressed as

$$-\sum_{m=1}^M \frac{C_m(\theta)}{\sum_{m=1}^M C_m(\theta)} \log \left(\frac{C_m(\theta)}{\sum_{m=1}^M C_m(\theta)} \right). \quad (7)$$

Definition 4.1 and 4.2 are also considered in in (Li et al., 2020) and offer metric definitions for evaluating performance equity among agents. Specifically, higher variance or a lower $\mathbb{H}_{\text{norm}}(C(\theta))$ indicates larger variability (*i.e.*, less equity) in the performance across agents.

We next provide theorems showing that the Equitable Objective Eq. (1) can encourage a more fair solution according to Definition 4.1 and 4.2. We initiate the analysis with the special case of $q = 1$, and prove that $q = 1$ can lead to a more equitable performance distribution than $q = 0$. The notation θ_q^* denotes the global optimal solution of $\min_{\theta} \mathcal{L}^q(\theta)$.

Theorem 4.3. *When $q = 1$, the optimum of Equitable Objective is more equitable compared to $q = 0$, indicated by smaller variance of the model performance distribution, *i.e.* $\text{Var}(C_1(\theta_{q=1}^*), \dots, C_M(\theta_{q=1}^*)) < \text{Var}(C_1(\theta_{q=0}^*), \dots, C_M(\theta_{q=0}^*))$.*

Moving forward to the general case, we show that for any $q > 0$, the proposed Equitable Objective can achieve better uniformity in performance distribution given a small increase of q .

Theorem 4.4. *Let $C(\theta)$ be twice differentiable in θ with $\nabla^2 C(\theta) > 0$ (positive definite), for any $M \in \mathbb{N}$, the derivative of $\mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))$ w.r.t. the evaluation point p is non-negative, *i.e.*,*

$$\frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \Big|_{p=q} \geq 0. \quad (8)$$

Theorem 4.4 establishes that a positive partial derivative of $\mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))$ signifies that a small increase in p is associated with a greater degree of performance uniformity in the learning outcome (Beirami et al., 2019).

4.2. Generalization Bounds

Denote h as the hypothesis function of the public model, *i.e.* $h(x) = f(x, \theta)$. In this work, we prove that the proposed Equitable Objective in Eq. (2) enables the public model to generalize well on the equitable loss described in Eq. (9) (Mohri et al., 2019).

$$\mathcal{J}_{\kappa}(h) = \sum_{m=1}^M \kappa_m \mathbb{E}_{(x,y) \sim D_m} C_m(h(x), y), \quad (9)$$

where $\kappa = [\kappa_1, \dots, \kappa_M]$ lies in a probability simplex Δ .

To show the generalization bound, we first give an equivalence of the Equitable Objective in Eq. (2). Given the definition of dual norm, we have

$$\begin{aligned} \tilde{\mathcal{L}}^q(h) &= (\mathcal{L}^q(h))^{\frac{1}{q+1}} \\ &= \max_{v, \|v\|_p \leq 1} \sum_{m=1}^M \left(\frac{v_j}{N_m} \sum_{i=1}^{N_m} C_m(h(x_{m,i}), y_{m,i}) \right), \end{aligned} \quad (10)$$

where $\frac{1}{p} + \frac{1}{q+1} = 1$ ($p \geq 1, q \geq 0$). Thus, the proposed Equitable Objective in Eq. (2) is equivalent to minimizing the empirical loss $\tilde{\mathcal{L}}^q(h)$ in Eq. (10). We present the generalization bound for $\mathcal{J}_{\kappa}(h)$ which depends on $\tilde{\mathcal{L}}^q(h)$ as below.

Proposition 4.5. *Assume that the cost functions cost_m are bounded by B . Then for any $\delta > 0$, with probability at least $1 - \delta$, the following holds for any κ in a probability simplex Δ , and any $h \in H$:*

$$\begin{aligned} \mathcal{J}_{\kappa}(h) &\leq \max_{\kappa \in \Delta} (\|\kappa\|_p) \tilde{\mathcal{L}}^q(h) + \max_{\kappa \in \Delta} \left(\mathbb{E}[\max_{h \in H} \mathcal{J}_{\kappa}(h) \right. \\ &\quad \left. - \mathcal{L}_{\kappa}(h)] \right) + B \left(\sqrt{\sum_m \frac{\kappa_m^2}{2N_m} \log \frac{1}{\delta}} \right), \end{aligned} \quad (11)$$

where $\frac{1}{p} + \frac{1}{q+1} = 1$, $\tilde{\mathcal{L}}^q(h)$ is the equivalent Equitable Objective in Eq. (10), and $\mathcal{L}_{\kappa}(h) = \sum_{m=1}^M \frac{\kappa_m}{N_m} \sum_{i=1}^{N_m} C_m(h(x_{m,i}), y_{m,i})$ is the empirical loss of $\mathcal{J}_{\kappa}(h)$.

5. Empirical Case Studies

We evaluate the effectiveness of EQUITABLE PM in fostering a more equitable solution for downstream heterogeneous agents, each with their own business objective, while utilizing the prediction from an upstream public model. Our empirical study encompasses the applications of data centers and Electric Vehicles (EV) charging.

Evaluation Metrics Instead of solely prioritizing the prediction accuracy, we emphasize the outcome, *e.g.*, decision cost (or rewards), of downstream agents from using the prediction of the upstream public model. Moreover, for diverse agents, we believe the algorithm should promote an equitable/uniform distribution of performance rather than disproportionately affecting specific agents. Our evaluation therefore incorporates three key metrics to assess the uniformity of performance distribution across agents: 1) *Variance* of the cost regret; 2) *Mean* of the cost regret; and 3) $C_{95} - C_5$ *percentile*, the discrepancy between the 95% and 5% percentiles of the cost regret across agents.

5.1. Application I: Carbon Efficiency in Data Centers

Setup Data centers are responsible for a significant amount of energy consumption and carbon emissions. In order to reduce their carbon footprint, it is crucial to manage energy consumption and optimize the allocation of workloads (Radovanović et al., 2022; Patterson et al., 2022).

In empirical studies, we denote the workload demand of data center j at time step t as $w_{j,t}$, represent the allocated computational resource as $p_{j,t}$, and indicate the predicted carbon emission rate at time t by c_t , where c_t is estimated by a public model. The processing delay can be calculated as $\frac{w_{j,t}}{p_{j,t} - w_{j,t}}$. Our objective is to minimize the combined impact of carbon emissions, $p_{j,t}c_t$, and processing latency, $\frac{w_{j,t}}{p_{j,t} - w_{j,t}}$, by determining the optimal allocation of computational resource p_t , as shown in Eq. (12),

$$\min_{p_{j,t}} p_{j,t}c_t + \lambda_j \frac{w_{j,t}}{p_{j,t} - w_{j,t}}, \quad (12)$$

where λ_j adjusts the relative significance of carbon emissions and processing latency for different data centers.

Datasets Our experiments mainly use the publicly available state-level energy fuel mix dataset (U.S. Energy Information Administration) and the Azure cloud workload dataset (Shahrad et al., 2020). The fuel mix dataset provides information on various energy sources utilized in electricity generation (e.g. coal, natural gas, and oil) while the Azure cloud workload dataset captures the energy consumption/demand patterns of the cloud center across different time periods. Besides, we utilize the carbon conversion rates provided in (Gao et al., 2012) to calculate the carbon emissions associated with different types of fuel used for energy generation. More details are in Appendix A.3.2.

Implementation Details We set the number of downstream agents as 50. We set up 3 different settings by varying data distribution and the values of λ among agents. The 50 agents have Wasserstein distance of w_j ranges falling within $[0.03, 0.58]$ and they are labeled as “similar agents”. At the same time, we randomly select 20 agents from the set and introduce random noise, resulting the total 50 agents with Wasserstein distance w.r.t. w_j spanning $[0.04, 57.97]$, which are labeled as “different agents”. Additionally, regarding the values of λ , “same λ ” designates $\lambda = 2$, whereas “different λ ” spans $\lambda = \{2, 4, \dots, 100\}$ among agents. Given the time-series nature of the datasets, we train and employ an LSTM network as the shared public model. More details are provided in Appendix A.3.1.

Results In Table 1, we present the performance comparison between EQUITABLE PM and the traditional public model that does not consider the decision-making process

of downstream agents, referred to as the Plain PM. From the Table 1, we can observe that the values of cost regret variance and $C_{95} - C_5$ achieved by EQUITABLE PM are smaller than the Plain PM. The variance and percentile measure values of EQUITABLE PM decrease as the value of q increases, suggesting more uniform cost regret distributions, and therefore a fairer solution. Also, the EQUITABLE PM has resulted in an improved cost regret mean compared to the Plain PM. Although the EQUITABLE PM does not achieve the minimum MSE on predicting carbon emissions c_t , it delivers more equitable and accurate cost outcomes for heterogeneous downstream agents under various settings.

Figure 2 shows the distribution of cost regret with various q values under different setups. When the data distribution among agents remains similar but with varying λ values, an increase in the value of q leads to a distribution with lower variance, as observed in Figure 2 (a). In Figure 2 (b), (c) and (d), we observe that the cost regret distributions become less dispersed when q increases, indicating a more equitable solution for different agents.

5.2. Application II: Scheduled EV Charging for Environmental Sustainability

Setup The increasing popularity of EV raises concerns about their environmental impact. To address this, scheduling EV charging can play a pivotal role in enhancing both environmental sustainability and the stability of the power system (Filote et al., 2020). Here, we evaluate the potential of EQUITABLE PM for a more equitable solution, in the context of optimizing the EV charging schedule aiming at minimizing the financial cost, together with carbon emission and water consumption.

Consider an EV j with an initial electrical charge state, denoted as I_j , which requires attaining an electric charge level represented as D_j . This charging process occurs within a defined time window that begins at s_j and concludes at e_j . For optimization purpose, we discretize the time window $[s_j, e_j]$ into time slots $\tau = \{1, \dots, T\}$ and utilize a binary charging schedule defined as X_j . In the schedule, each element $x_{j,t}$ is either 1, indicating that we charge the vehicle at the time t , or 0 if we don’t (e.g., $X_j = [1, 0, \dots, 1]$, with a total of $|\tau|$ elements in X_j). And the amount of electricity charged at each time step t of the j -th EV is $\zeta_{j,t}$.

At each time step t within $[s_j, e_j]$, an upstream public model predicts the combined carbon, water efficiency and electricity price, expressed as $E_t = E_t^C + \gamma E_t^W + \eta E_t^P$, where E_t^C and E_t^W denote the carbon and water efficiency at time t , respectively, while E_t^P represents the electricity price at time t for downstream agent EV. The term γ and η represents their relative weight of these factors. Here, the efficiency refers to the amount of carbon emission or water consumption per

Table 1: Statistics of the test results under different setups. As q increases, the variance and $C_{95} - C_5$ percentile of cost regret distribution across agents decrease, suggesting a more uniform distribution of costs across groups. The EQUITABLE PM also achieves lower means in costs regrets across agents compared to the Plain PM in general.

Setting	Method	$q + 1$	Variance	Mean	$C_{95} - C_5$	MSE
Similar Agents, Different λ	EQUITABLE PM	1	0.0029	0.1591	0.1687	4.6308
		1.1	0.0003	0.0544	0.0576	4.2465
		1.5	0.0002	0.0465	0.0493	4.2194
	Plain PM	-	0.0085	0.2732	0.2897	4.2054
Different Agents, Same λ	EQUITABLE PM	1	0.0008	0.0909	0.0809	5.0991
		3	0.0001	0.0345	0.0306	4.4338
		20	1.71e-5	0.0136	0.0121	4.2028
	Plain PM	-	0.0009	0.0988	0.0879	4.2013
Different Agents, Different λ	EQUITABLE PM	1	0.0181	0.2619	0.4229	4.6607
		3	0.0068	0.1603	0.2588	4.4182
		10	0.0055	0.1444	0.2331	4.3819
	Plain PM	-	0.0602	0.4779	0.7717	4.2013

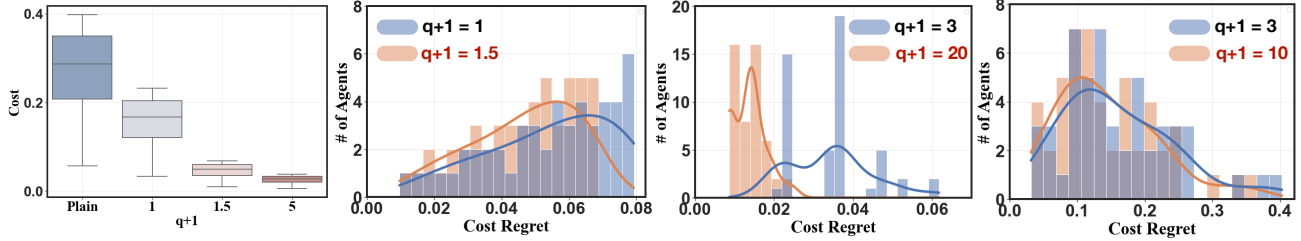


Figure 2: (a) Comparison of cost regret distributions between Plain PM vs. EQUITABLE PM on “similar agents” with different λ . The EQUITABLE PM shows lower variability in cost distribution compared to the Plain PM. With varied $q + 1$, we show cost regret distributions when using EQUITABLE PM in (b) “similar agents” with different λ ; (c) “different agents” with same λ ; (d) “different agents” with different λ . As the value of q increases, the cost regret distribution across downstream agents achieves greater uniformity, implying a more equitable solution.

unit of electricity generated.

The objective is to reduce the total cost, which includes carbon emissions, water consumption, and the financial cost of electricity incurred throughout the charging process, by determining the optimal charging schedule for the j -th EV. We formulate the objective in Eq. (13).

$$\begin{aligned}
 & \min_{X_j} \sum_t \zeta_{j,t} x_{j,t} \cdot E_t \\
 & s.t. \quad I_j + \sum_t \zeta_{j,t} x_{j,t} = D_j \\
 & \quad E_t = E_t^C + \gamma E_t^W + \eta E_t^P,
 \end{aligned} \tag{13}$$

where $X_j = [x_{j,1}, \dots, x_{j,T}]$.

Datasets Our main sources of datasets include the publicly available ACN-Data, collected from the Caltech ACN and similar websites (Lee et al., 2019), as well as the California Electricity Market (CAISO) (CAISO). The ACN-Data records the real time charging details, including EV arrival/departure times and actual energy delivered in each charging session. Simultaneously, the CAISO provides data on electricity prices in California. We use the ACN-Data to

estimate power demand and charging rates for EV in residential areas, considering that EV models are similar between residential and other charging stations (Wang & Paranjape, 2015). Additionally, we use the state-level energy fuel mix data (U.S. Energy Information Administration) for carbon and water efficiency calculation. Regarding the available charging time window, from s_j to e_j , in residential sectors, we use the data from The National Household Travel Survey (NHTS) to approximate (U.S. Department of Transportation, 2017; Wang & Paranjape, 2015). Further details can be founded in Appendix A.3.2.

Implementation Details In the experiments, we follow the calculation of carbon and water efficiency outlined in (Li et al., 2023a). We recognize that different EV exhibit distinct charging patterns (Sun et al., 2020). Given our central objective of ensuring fairness across a diverse range of EV, we here however opt for a simplifying assumption of a uniform charging rate, implying that $\zeta_{j,t}$ remain constant w.r.t. t for the j -th EV (Sun et al., 2020). More specifically, this rate is calculated by the charged electricity divided by the difference between the ending charging and starting times of the j -th EV, as provided in the ACN-Data. Additionally,

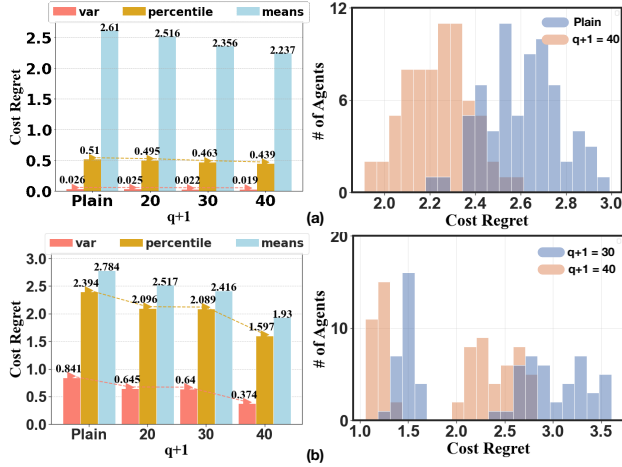


Figure 3: Statistics and cost regret distributions of test result between Plain PM and EQUITABLE PM with varied $q + 1$ for (a) “similar”; and (b) “different” agents. The EQUITABLE PM demonstrates improved uniformity in agent distributions compared to the Plain PM. As q increases, the uniformity of cost distribution across agents improves.

we use the energy demands of each EV provided in the ACN-Data as D_j . To ensure flexibility in charging scheduling, we set the time frame $|\tau| = 12$. For instance, if we use an hourly unit, this corresponds to scheduling charging for half of the day. We set the number of downstream EV agents as 70 and γ and η are set at 1. In the experiments, we explore the effectiveness of EQUITABLE PM across agents exhibiting varied data distributions. The Wasserstein distance range of $(D_j - I_j)$ for agents labeled as “similar agents” spans $[0.80, 9.00]$, while for “different agents”, it ranges $[1.33, 60.36]$. The *Transformer* architecture, with a linear layer as the task head, is employed as the shared public model to predict E_t for scheduling downstream EV charging. More details are provided in Appendix A.3.1.

Results Figure 3 reports the evaluation results between the Plain PM and EQUITABLE PM with different q for agents exhibiting both similar and different distributions. The results demonstrate that the EQUITABLE PM consistently achieves lower variance and $C_{95} - C_5$ percentile values compared to utilizing the Plain PM in both settings. Examining Figure 3, it becomes evident that as the value of q increases, both the variance and the range $C_{95} - C_5$ percentile of cost regret distributions among agents decrease, indicating a trend towards a more uniformly distributed performance.

6. Related Works

Fairness in Machine Learning Fairness is a prevalent topic within the realm of machine learning, often focusing on the protection of certain groups or attributes. The problem stems partly from inherent biases within datasets and

could be further magnified by models (Wan et al., 2023; Li et al., 2023b). Various approaches have been developed to mitigate this form of unfairness, spanning different stages of model development. These approaches encompass pre-processing methods, such as excluding sensitive attributes from the datasets to prevent model reliance on these factors (Biswas & Rajan, 2021; Madras et al., 2018a). Post-processing techniques calibrate prediction outcomes after training (Pessach & Shmueli, 2022; Noriega-Campero et al., 2018), and in-processing methodologies directly integrates fairness considerations during model training (Wan et al., 2023; Kearns et al., 2018).

Our work enforces fairness during training but takes a distinct perspective. We emphasize the equity/uniformity of performance distribution across heterogeneous agents, as we view the upstream public model as a shared resource serving diverse downstream agents. While certain studies advocate for equivalent error rates as a fairness criterion (Cotter et al., 2019), our goal does not prioritize optimizing equal model accuracy across all agents. Drawing an analogy between the shared public model and a resource, we are inspired by a unified resource allocation framework called α -fairness, where the service provider can adjust fairness emphasis via a single hyperparameter (Mo & Walrand, 2000; Lan et al., 2009). However, the aspect of equity, specifically concerning the impact of predictions from a shared public model on diverse downstream agents’ business decisions, a focal point in our work, remains unexplored in previous literatures.

Decision-focused Learning Decision-focused learning is an emerging area in machine learning that trains a model to optimize decisions by integrating prediction and optimization within an end-to-end system (Mandi et al., 2023). It diverges from the predict-then-optimize framework (Balghithi et al., 2022; Elmachetoub & Grigas, 2020), where a ML model is trained initially to map observed features to relevant parameters of a combinatorial optimization problem, followed by using a specialized optimization algorithm to solve the decision problem based on predicted parameters. The predict-then-optimize methodology assumes accurate predictions generate precise models, enabling optimal decisions. However, ML models often lack perfect accuracy, prediction errors thus can lead to suboptimal decisions.

In comparison, decision-focused learning directly trains the ML model to make predictions that lead to good decisions, where the optimization is embedded as a component of the ML model, creating an end-to-end approach. Recent studies have utilized supervised or reinforcement learning to optimize ultimate decisions with end-to-end machine learning (Wilder et al., 2020; Johnson-Yu et al., 2023; Bello et al., 2017; Donti et al., 2019). This holistic approach has enhanced the model’s capability to drive informed and effective downstream decisions. However, few existing works

have considered the issue of performance disparity across diverse business agents, each with their distinct concerns, specifically in the context of using a publicly shared model to optimize their decisions (Yang et al., 2023; Madras et al., 2018b; Wilder et al., 2021).

7. Conclusion

In this paper, we introduce the novel Equitable Objective and its corresponding solver, the EQUITABLE PM with either differentiable or non-differentiable cost functions, to promote the performance equity/uniformity among diverse downstream agents that depend on the predictions of a shared public model for their decision-making. Alongside theoretical proofs demonstrating the performance uniformity improvement achieved by our proposed approach, the empirical case studies using real-world datasets further validates that EQUITABLE PM can attain a more equitable solution compared to methods that solely focuses on minimizing the prediction error without considering the objectives of downstream agents in different settings.

Limitation & Future Works Our current method relies on accessing the decision costs from downstream groups to construct a socially-responsible public model, potentially raising privacy and security concerns. In future research, we aim to investigate ways that uphold privacy and increase robustness against adversarial attacks (e.g., maliciously reporting decision costs) when extending our approach. Furthermore, while the models used in the current case studies are appropriate for the present context, their scale is relatively modest, also due to a constraint imposed by our limited computing resources. We would like to explore the efficacy of our proposed method in more extensive architectures and other domains such as healthcare. Additionally, the exploration of alternative methods, such as using fine-tuning to align public foundation models for making business-informed decisions and addressing fairness concerns accordingly, continues to be a key focus for upcoming research endeavors.

Acknowledgements

We would like to first thank the anonymous reviewers for their insightful comments.

Yejia Liu, Jianyi Yang, Pengfei Li, and Shaolei Ren were supported in part by the US NSF under grants CNS-1910208, CNS-2007115, and CCF-2324941.

Tongxin Li was supported in part by the National Natural Science Foundation of China (NSFC) under grant No. 72301234, Pengcheng Peacock Research Fund (Category C), the Guangdong Key Lab of Mathematical Foundations for AI (2023B1212010001), the Shenzhen Key Lab of Crowd

Intelligence Empowered Low-Carbon Energy Network, and the start-up funding UDF01002773 of CUHK-Shenzhen.

Impact Statement

This paper presents work whose goal is to advance the field of Machine Learning to make public AI more equitable when serving multiple agents each having distinct downstream decision processes and objectives. There are many potential societal consequences of our work. Notably, our work can lead to more uniform decision costs among multiple agents sharing a single public model.

References

- Agarwal, A., Jin, Y., and Zhang, T. *Voq1: Towards optimal regret in model-free rl with nonlinear function approximation*. In *The Thirty Sixth Annual Conference on Learning Theory*, pp. 987–1063. PMLR, 2023.
- Altman, E., Avrachenkov, K., and Garnaev, A. Generalized α -fair resource allocation in wireless networks. In *2008 47th IEEE Conference on Decision and Control*, pp. 2414–2419, 2008. doi: 10.1109/CDC.2008.4738709.
- Balghiti, O. E., Elmachtoub, A. N., Grigas, P., and Tewari, A. Generalization bounds in the predict-then-optimize framework, 2022.
- Barocas, S., Hardt, M., and Narayanan, A. *Fairness and machine learning: Limitations and opportunities*. MIT Press, 2023.
- Beirami, A., Calderbank, R., Christiansen, M. M., Duffy, K. R., and Medard, M. A characterization of guesswork on swiftly tilting curves. *IEEE Transactions on Information Theory*, 65(5):2850–2871, May 2019. ISSN 1557-9654. doi: 10.1109/tit.2018.2879477. URL <http://dx.doi.org/10.1109/TIT.2018.2879477>.
- Bello, I., Pham, H., Le, Q. V., Norouzi, M., and Bengio, S. Neural combinatorial optimization with reinforcement learning, 2017.
- Bhandari, J. and Russo, D. Global optimality guarantees for policy gradient methods. *Operations Research*, 2024.
- Biswas, S. and Rajan, H. Fair preprocessing: towards understanding compositional fairness of data transformers in machine learning pipeline. In *Proceedings of the 29th ACM Joint Meeting on European Software Engineering Conference and Symposium on the Foundations of Software Engineering, ESEC/FSE ’21*. ACM, August 2021. doi: 10.1145/3468264.3468536. URL <http://dx.doi.org/10.1145/3468264.3468536>.

- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosse-
lut, A., Brunskill, E., et al. On the opportunities and risks
of foundation models. *arXiv preprint arXiv:2108.07258*,
2021.
- CAISO. California iso - market price maps. URL
[http://www.caiso.com/pricemap/Pages/
default.aspx](http://www.caiso.com/pricemap/Pages/default.aspx).
- Cotter, A., Jiang, H., Wang, S., Narayan, T., You, S., Srid-
haran, K., and Gupta, M. R. Optimization with non-
differentiable constraints with applications to fairness,
recall, churn, and other goals. *Journal of Machine Learn-
ing Research*, 2019.
- Donti, P. L., Amos, B., and Kolter, J. Z. Task-based end-to-
end model learning in stochastic optimization, 2019.
- Elmachtoub, A. N. and Grigas, P. Smart "predict, then
optimize", 2020.
- Filote, C., Felseghi, R.-A., Raboaca, M. S., and Aşchilean,
I. Environmental impact assessment of green energy
systems for power supply of electric vehicle charging
station. *International Journal of Energy Research*, 44
(13):10471–10494, 2020.
- Gao, P. X., Curtis, A. R., Wong, B., and Keshav, S. It's
not easy being green. In *Proceedings of the ACM
SIGCOMM 2012 Conference on Applications, Tech-
nologies, Architectures, and Protocols for Computer
Communication*, SIGCOMM '12, pp. 211–222, New
York, NY, USA, 2012. Association for Computing
Machinery. ISBN 9781450314190. doi: 10.1145/
2342356.2342398. URL [https://doi.org/10.
1145/2342356.2342398](https://doi.org/10.1145/2342356.2342398).
- Jang, J. and Yang, H. J. α -fairness-maximizing user asso-
ciation in energy-constrained small cell networks. *IEEE
Transactions on Wireless Communications*, 21(9):7443–
7459, 2022. doi: 10.1109/TWC.2022.3158694.
- Johnson-Yu, S., Wang, K., Finocchiario, J., Taneja, A.,
and Tambe, M. *Modeling Robustness in Decision-
Focused Learning as a Stackelberg Game*, pp. 2908–2909.
International Foundation for Autonomous Agents and
Multiagent Systems, Richland, SC, 2023. ISBN
9781450394321.
- Kearns, M., Neel, S., Roth, A., and Wu, Z. S. Prevent-
ing fairness gerrymandering: Auditing and learning for
subgroup fairness, 2018.
- Lan, T., Kao, D., Chiang, M., and Sabharwal, A. An ax-
iomatic theory of fairness in network resource allocation,
2009.
- Lee, Z., Li, T., and Low, S. Acn-data: Analysis and appli-
cations of an open ev charging dataset. pp. 139–149, 06
2019. doi: 10.1145/3307772.3328313.
- Li, P., Yang, J., Islam, M. A., and Ren, S. Making ai less
"thirsty": Uncovering and addressing the secret water
footprint of ai models, 2023a.
- Li, T., Sanjabi, M., Beirami, A., and Smith, V. Fair resource
allocation in federated learning, 2020.
- Li, T., Guo, Q., Liu, A., Du, M., Li, Z., and Liu, Y. Fairer:
Fairness as decision rationale alignment, 2023b.
- Madras, D., Creager, E., Pitassi, T., and Zemel, R. Learning
adversarially fair and transferable representations, 2018a.
- Madras, D., Pitassi, T., and Zemel, R. Predict responsi-
bly: improving fairness and accuracy by learning to defer.
In *Proceedings of the 32nd International Conference on
Neural Information Processing Systems*, NIPS'18, pp.
6150–6160, Red Hook, NY, USA, 2018b. Curran Asso-
ciates Inc.
- Malik, D., Pananjady, A., Bhatia, K., Khamaru, K., Bartlett,
P., and Wainwright, M. Derivative-free methods for policy
optimization: Guarantees for linear quadratic systems. In
Chaudhuri, K. and Sugiyama, M. (eds.), *Proceedings of
the Twenty-Second International Conference on Artificial
Intelligence and Statistics*, volume 89 of *Proceedings of
Machine Learning Research*, pp. 2916–2925. PMLR, 16–
18 Apr 2019. URL [https://proceedings.mlr.
press/v89/malik19a.html](https://proceedings.mlr.press/v89/malik19a.html).
- Mandi, J., Kotary, J., Berden, S., Mulamba, M., Bucarey,
V., Guns, T., and Fioretto, F. Decision-focused learn-
ing: Foundations, state of the art, benchmark and future
opportunities, 2023.
- Mo, J. and Walrand, J. Fair end-to-end window-based con-
gestion control. *IEEE/ACM Transactions on Networking*,
8(5):556–567, 2000. doi: 10.1109/90.879343.
- Moerland, T. M., Broekens, J., Plaat, A., Jonker, C. M.,
et al. Model-based reinforcement learning: A survey.
Foundations and Trends® in Machine Learning, 16(1):
1–118, 2023.
- Mohri, M., Sivek, G., and Suresh, A. T. Agnostic federated
learning, 2019.
- Nguyen, T., Brandstetter, J., Kapoor, A., Gupta, J. K., and
Grover, A. Climax: A foundation model for weather and
climate. *arXiv preprint arXiv:2301.10343*, 2023.
- Noriega-Campero, A., Bakker, M. A., Garcia-Bulle, B.,
and Pentland, A. Active fairness in algorithmic decision
making, 2018.

- Patterson, D., Gonzalez, J., Hölzle, U., Le, Q., Liang, C., Munguia, L.-M., Rothchild, D., So, D. R., Texier, M., and Dean, J. The carbon footprint of machine learning training will plateau, then shrink. *Computer*, 55(7):18–28, 2022.
- Pessach, D. and Shmueli, E. A review on fairness in machine learning. *ACM Computing Surveys (CSUR)*, 55(3):1–44, 2022.
- Peters, J. and Schaal, S. Policy gradient methods for robotics. In *2006 IEEE/RSJ International Conference on Intelligent Robots and Systems*, pp. 2219–2225. IEEE, 2006.
- Radovanović, A., Koningstein, R., Schneider, I., Chen, B., Duarte, A., Roy, B., Xiao, D., Haridasan, M., Hung, P., Care, N., et al. Carbon-aware computing for datacenters. *IEEE Transactions on Power Systems*, 38(2):1270–1280, 2022.
- Shah, D., Osinski, B., Ichter, B., and Levine, S. Lm-nav: Robotic navigation with large pre-trained models of language, vision, and action, 2022.
- Shahrad, M., Fonseca, R., Íñigo Goiri, Chaudhry, G., Batum, P., Cooke, J., Laureano, E., Tresness, C., Russinovich, M., and Bianchini, R. Serverless in the wild: Characterizing and optimizing the serverless workload at a large cloud provider, 2020.
- Smart, J. G. and Salisbury, S. D. Plugged in: How americans charge their electric vehicles. 7 2015. doi: 10.2172/1369632. URL <https://www.osti.gov/biblio/1369632>.
- Sun, C., Li, T., Low, S., and Li, V. Classification of electric vehicle charging time series with selective clustering. *Electric Power Systems Research*, 189:106695, 12 2020. doi: 10.1016/j.epsr.2020.106695.
- Sun, J. J., Kennedy, A., Zhan, E., Anderson, D. J., Yue, Y., and Perona, P. Task programming: Learning data efficient behavior representations. In *2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pp. 2875–2884, 2021. doi: 10.1109/CVPR46437.2021.00290.
- U.S. Department of Transportation. Federal Highway Administration. National Household Travel Survey, 2017. URL <http://nhts.ornl.gov>.
- U.S. Energy Information Administration. Electricity – consumption of fuels used to generate electricity. URL <https://www.eia.gov/electricity/data.php>.
- Wan, M., Zha, D., Liu, N., and Zou, N. In-processing modeling techniques for machine learning fairness: A survey. *ACM Trans. Knowl. Discov. Data*, 17(3), mar 2023. ISSN 1556-4681. doi: 10.1145/3551390. URL <https://doi.org/10.1145/3551390>.
- Wang, Z. and Paranjape, R. Optimal scheduling algorithm for charging electric vehicle in a residential sector under demand response. In *2015 IEEE Electrical Power and Energy Conference (EPEC)*, pp. 45–49, 2015. doi: 10.1109/EPEC.2015.7379925.
- Wilder, B., Dilkina, B., and Tambe, M. Melding the data-decisions pipeline: Decision-focused learning for combinatorial optimization. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, pp. 1658–1665, 2019.
- Wilder, B., Ewing, E., Dilkina, B., and Tambe, M. End to end learning and optimization on graphs, 2020.
- Wilder, B., Horvitz, E., and Kamar, E. Learning to complement humans. In *Proceedings of the Twenty-Ninth International Joint Conference on Artificial Intelligence, IJCAI’20*, 2021. ISBN 9780999241165.
- Yang, S., Nachum, O., Du, Y., Wei, J., Abbeel, P., and Schuurmans, D. Foundation models for decision making: Problems, methods, and opportunities, 2023.
- Yu, T., Thomas, G., Yu, L., Ermon, S., Zou, J. Y., Levine, S., Finn, C., and Ma, T. Mopo: Model-based offline policy optimization. *Advances in Neural Information Processing Systems*, 33:14129–14142, 2020.

A. Appendix

In the appendix, we offer additional details to complement the main text. The content is organized as follows:

- Section A.1. Detailed calculations to derive the gradient in Eq. (3) of Section 3.2.
- Section A.2. Providing proofs for the theorems and propositions in Section 4.
- Section A.3. Additional empirical details of implementation, datasets, and results for case studies in Section 5.
- Section A.4. Additional experiments where downstream agents have different objective cost functions.
- Section A.5. Proposal of a combined objective that explicitly incorporates the loss of the public model, $\mathcal{L}_f$, via a balancing hyperparameter β , providing a more nuanced control over the tradeoff between fairness and accuracy. We also provide the empirical results associated with this objective.

A.1. Details of Computing the Gradient

Completing Section 3.2, we present a detailed derivation of $\nabla_{\theta}\mathbb{E}[\mathcal{L}^q(\theta)]$. By the definition of $\mathcal{L}^q(\theta)$ in the main context, the expected cost is defined as

$$\begin{aligned}\mathbb{E}[\mathcal{L}^q(\theta)] &= \mathbb{E}_{(X,Y,\hat{Y},\Xi)\sim p_{\theta}} \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] \\ &= \int p_{\theta}(X, Y, \hat{Y}, \Xi) \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] d(X, Y, \hat{Y}, \Xi).\end{aligned}\tag{14}$$

Subsequently, we obtain

$$\begin{aligned}\nabla_{\theta}\mathbb{E}[\mathcal{L}^q(\theta)] &= \int \nabla_{\theta} p_{\theta}(X, Y, \hat{Y}, \Xi) \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] d(X, Y, \hat{Y}, \Xi) \\ &= \int p_{\theta}(X, Y, \hat{Y}, \Xi) \nabla_{\theta} \log p_{\theta}(X, Y, \hat{Y}, \Xi) \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] d(X, Y, \hat{Y}, \Xi).\end{aligned}$$

By decomposing the joint distribution $p_{\theta}(X, Y, \hat{Y}, \Xi)$ based on the chain rule as

$$p_{\theta}(X, Y, \hat{Y}, \Xi) = P(\Xi) \cdot P(Y | X) \cdot \sigma_{\theta}(\hat{Y} | X) \cdot P(X),$$

we have

$$\begin{aligned}\nabla_{\theta}\mathbb{E}[\mathcal{L}^q(\theta)] &= \int p_{\theta}(X, Y, \hat{Y}, \Xi) \nabla_{\theta} \log [P(\Xi) \cdot P(Y | X) \cdot \sigma_{\theta}(\hat{Y} | X) \cdot P(X)] \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] d(X, Y, \hat{Y}, \Xi) \\ &= \int p_{\theta}(X, Y, \hat{Y}, \Xi) \nabla_{\theta} \log \sigma_{\theta}(\hat{Y} | X) \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] d(X, Y, \hat{Y}, \Xi) \\ &= \mathbb{E}_{(X,Y,\hat{Y},\Xi)\sim p_{\theta}} \left\{ \nabla_{\theta} \log \sigma_{\theta}(\hat{Y} | X) \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] \right\} \\ &= \mathbb{E}_{(X,Y,\hat{Y},\Xi)\sim p_{\theta}} \left\{ \left[\sum_{m=1}^M \sum_{i=1}^{B_m} \nabla_{\theta} \log \sigma_{\theta}(\hat{y}_{m,i} | x_{m,i}) \right] \left[\sum_{m=1}^M \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,i} \right)^{q+1} \right] \right\}.\end{aligned}\tag{15}$$

Rewriting Eq. (15), we obtain the gradient stated in the Eq. (3).

A.2. Theoretical Proofs

A.2.1. PROOF OF THEOREM 4.3

Proof. Let $\theta_{q=0}^*$ and $\theta_{q=1}^*$ denote optimal solutions of $\min_{\theta} \mathcal{L}^{q=0}(\theta)$ and $\min_{\theta} \mathcal{L}^{q=1}(\theta)$ respectively. It follows that

$$\begin{aligned} \text{Var}(C_1(\theta_{q=1}^*), \dots, C_M(\theta_{q=1}^*)) &= \frac{1}{M} \sum_{m=1}^M C_m^2(\theta_{q=1}^*) - \left(\frac{1}{M} \sum_{m=1}^M C_m(\theta_{q=1}^*) \right)^2 \\ &\leq \frac{1}{M} \sum_{m=1}^M C_m^2(\theta_{q=0}^*) - \left(\frac{1}{M} \sum_{m=1}^M C_m(\theta_{q=1}^*) \right)^2 \\ &\leq \frac{1}{M} \sum_{m=1}^M C_m^2(\theta_{q=0}^*) - \left(\frac{1}{M} \sum_{m=1}^M C_m(\theta_{q=0}^*) \right)^2 = \text{Var}(C_1(\theta_{q=0}^*), \dots, C_M(\theta_{q=0}^*)), \end{aligned} \quad (16)$$

where the first inequality holds since $\theta_{q=1}^*$ minimizes $\frac{1}{M} \sum_{m=1}^M C_m^2(\theta_{q=1}^*)$, and the second inequality holds because $\theta_{q=0}^*$ minimizes $\frac{1}{M} \sum_{m=1}^M C_m(\theta_{q=0}^*)$. $\square$

A.2.2. PROOF OF THEOREM 4.4

To prove Theorem 4.4, it suffices to show that for any $q \in \mathbb{R}^+$, $M \in \mathbb{N}$, a small increase in q can result in a more equitable solution for the Equitable Objective, based on Definition 4.2. Specifically, we prove the derivative of $H_{\text{norma}}(C^{q+1}(\theta_p^*))$ w.r.t. the variable p at the point $p = q$ is non-negative, i.e.,

$$\frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \Big|_{p=q} \geq 0. \quad (17)$$

Proof of the statement above. For simplicity of notation, we denote the gradient of $C^{q+1}(\theta)$ with respect to θ as the vector $\nabla_{\theta} C^{q+1}(\theta)$, and the second order derivative of $C^{q+1}(\theta)$ with respect to θ as the Hessian matrix $\nabla_{\theta}^2 C^{q+1}(\theta)$. If $C(\theta) \neq 0$, we can easily verify that the Hessian matrix $\nabla^2 C^{q+1}(\theta)$ is positive definite for all $q \geq 0$. More specifically, we have

$$\nabla_{\theta} (\nabla_{\theta} C^{q+1}(\theta)) = (q+1) \nabla_{\theta} (C^q(\theta) \nabla_{\theta} C(\theta)) = (q+1) C^q(\theta) \nabla_{\theta}^2 C(\theta) + (q+1) q C^{q-1}(\theta) \nabla_{\theta} C(\theta) \nabla_{\theta} C(\theta)^{\top}. \quad (18)$$

By definition, $\nabla_{\theta}^2 C(\theta)$ is positive definite and $C(\theta) \nabla_{\theta} C(\theta)^{\top}$ is semi-positive definite. Since all the coefficients are non-negative, we conclude the Hessian matrix $\nabla_{\theta}^2 C^q(\theta)$ is positive definite when $C(\theta) \neq 0$.

If $C(\theta) = 0$, both the vector $\nabla_{\theta} C^{q+1}(\theta)$ and the matrix $\nabla_{\theta}^2 C^{q+1}(\theta)$ are equal to zero.

Subsequently, the proof of Eq. (17) proceeds as follows:

$$\begin{aligned} \frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \Big|_{p=q} &= - \frac{\partial}{\partial p} \sum_m \frac{C_m^{q+1}(\theta_p^*)}{\sum_m C_m^{q+1}(\theta_p^*)} \ln \left(\frac{C_m^{q+1}(\theta_p^*)}{\sum_m C_m^{q+1}(\theta_p^*)} \right) \Big|_{p=q} \\ &= - \frac{\partial}{\partial p} \sum_m \frac{C_m^{q+1}(\theta_p^*)}{\sum_m C_m^{q+1}(\theta_p^*)} \ln \left(C_m^{q+1}(\theta_p^*) \right) \Big|_{p=q} + \frac{\partial}{\partial p} \ln \sum_m C_m^{q+1}(\theta_p^*) \Big|_{p=q}. \end{aligned} \quad (19)$$

For the second term in Eq. (19), we have

$$\begin{aligned} \frac{\partial}{\partial p} \ln \sum_m C_m^{q+1}(\theta_p^*) \Big|_{p=q} &= \frac{\sum_m \nabla_{\theta} C_m^{q+1}(\theta_p^*)^{\top} \cdot \frac{\partial \theta_p^*}{\partial p}}{\sum_m C_m^{q+1}(\theta_p^*)} \Big|_{p=q} \\ &= \frac{1}{\sum_m C_m^{q+1}(\theta_p^*)} \cdot \frac{\partial \theta_p^*}{\partial p} \Big|_{p=q} \cdot \sum_m \nabla_{\theta} C_m^{q+1}(\theta_p^*). \end{aligned} \quad (20)$$

Since the θ_p^* is an optimal solution for the $\mathcal{L}^p(\theta)$ objective, then for $q = p$, by definition we have $\sum_m \nabla_{\theta} C_m^{q+1}(\theta_p^*) = 0$.

Therefore, the second term of Eq. (21) is zero. The derivative can then be rewritten as

$$\begin{aligned}
 \frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \Big|_{p=q} &= - \sum_m \frac{(\frac{\partial}{\partial p} \theta_p^*|_{p=q})^\top \nabla_\theta C_m^{q+1}(\theta_p^*)}{\sum_m C_m^{q+1}(\theta_p^*)} \ln(C_m^{q+1}(\theta_p^*)) \\
 &\quad - \sum_m \frac{C_m^{q+1}(\theta_p^*)}{\sum_m C_m^{q+1}(\theta_p^*)} \frac{(\frac{\partial}{\partial p} \theta_p^*|_{p=q-1})^\top \nabla_\theta C_m^{q+1}(\theta_p^*)}{C_m^{q+1}(\theta_p^*)} \\
 &= - \sum_m \frac{(\frac{\partial}{\partial p} \theta_p^*|_{p=q})^\top \nabla_\theta C_m^{q+1}(\theta_p^*)}{\sum_m C_m^{q+1}(\theta_p^*)} (\ln(C_m^{q+1}(\theta_p^*)) + 1).
 \end{aligned} \tag{21}$$

Here, if for all $M \in \mathbb{N}$, the costs $C_m(\theta_p^*)$ are all zero costs. Therefore, we see that $\frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \Big|_{p=q} = 0$, leading to the desirable result.

For the non-trivial case, there exists some $M \in \mathbb{N}$ such that $C_m(\theta_p^*) > 0$. Since θ_p^* is an optimal solution of our objective function, we have $\sum_m \nabla_\theta C_m^{p+1}(\theta_p^*) = 0$ for all $p \geq 0$. In other words, $\frac{\partial}{\partial p} \sum_m \nabla_\theta C_m^{p+1}(\theta_p^*) = 0$. Then we can calculate the gradient as follows

$$\begin{aligned}
 &\frac{\partial}{\partial p} \sum_m \nabla_\theta C_m^{p+1}(\theta_p^*) \\
 &= \sum_m \nabla_\theta^2 C_m^{p+1}(\theta_p^*) \frac{\partial}{\partial p} \theta_p^* + \sum_m \left(C_m^p(\theta_p^*) + (p+1)C_m^p(\theta_p^*) \ln(C_m(\theta_p^*)) \right) \nabla_\theta C_m(\theta_p^*) \\
 &= \sum_m \nabla_\theta^2 C_m^{p+1}(\theta_p^*) \frac{\partial}{\partial p} \theta_p^* + \frac{1}{p+1} \sum_m \left((p+1)C_m^p(\theta_p^*) \nabla_\theta C_m(\theta_p^*) \right) + \left((p+1)C_m^p(\theta_p^*) \nabla_\theta C_m(\theta_p^*) \ln(C_m(\theta_p^*)) \right) \\
 &= \sum_m \nabla_\theta^2 C_m^{p+1}(\theta_p^*) \frac{\partial}{\partial p} \theta_p^* + \frac{1}{p+1} \sum_m \left(\ln(C_m^{p+1}(\theta_p^*)) + 1 \right) \nabla_\theta C_m^{p+1}(\theta_p^*).
 \end{aligned} \tag{22}$$

To summarize, we have

$$\sum_m \nabla_\theta^2 C_m^{p+1}(\theta_p^*) \frac{\partial}{\partial p} \theta_p^* + \frac{1}{p+1} \sum_m (\ln(C_m^{p+1}(\theta_p^*)) + 1) \nabla_\theta C_m^{p+1}(\theta_p^*) = 0. \tag{23}$$

In our non-trivial case, there exists at least one $m \in \mathbb{N}$, such that the Hessian matrix $\nabla_\theta^2 C_m^{p+1}(\theta_p^*)$ is positive definite. Then the matrix $\left(\sum_m \nabla_\theta^2 C_m^{p+1}(\theta_p^*) \right)$ is also positive definite. Therefore, we can calculate the gradient $\frac{\partial}{\partial p} \theta_p^*$ as below

$$\frac{\partial}{\partial p} \theta_p^* = - \frac{1}{p+1} \left(\sum_m \nabla_\theta^2 C_m^{p+1}(\theta_p^*) \right)^{-1} \sum_m (\ln(C_m^{p+1}(\theta_p^*)) + 1) \nabla_\theta C_m^{p+1}(\theta_p^*). \tag{24}$$

Plugging Eq. (24) into Eq. (21), we have

$$\begin{aligned}
 \frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \Big|_{p=q} &= \frac{\sum_m (\ln(C_m^{q+1}(\theta_p^*)) + 1) \nabla_\theta C_m^{q+1}(\theta_p^*)^\top}{(p+1) \sum_m C_m^{q+1}(\theta_p^*)} \left(\sum_m \nabla_\theta^2 C_m^{p+1}(\theta_p^*) \right)^{-1} \\
 &\quad \cdot \sum_m \nabla_\theta C_m^{p+1}(\theta_p^*) (\ln(C_m^{p+1}(\theta_p^*)) + 1) \Big|_{p=q}.
 \end{aligned} \tag{25}$$

Since the matrix $\left(\sum_m \nabla_\theta^2 C_m^q(\theta_p^*) \right)$ is positive definite and the coefficient $q \sum_m C_m^q(\theta_p^*)$ is positive, we conclude that

$$\frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \Big|_{p=q} \geq 0. \quad \square$$

As a result, Eq. (17) implies that for any p , the performance distribution of $\{C_1^p(\theta_{p+\epsilon}^*), \dots, C_M^p(\theta_{p+\epsilon}^*)\}$ exhibits greater uniformity compared to the distribution of $\{C_1^p(\theta_p^*), \dots, C_M^p(\theta_p^*)\}$, provided that the value of ϵ is sufficiently small.

Corollary A.1. Let $C(\theta)$ be twice differentiable in θ with $\nabla^2 C(\theta) > 0$ (positive definite), for the special case $M = 2$, the derivative of $\mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))$ w.r.t. the evaluation point p is non-negative for all $p \geq 0$ and $q \geq 0$, i.e.,

$$\frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \geq 0. \quad (26)$$

Proof. Let $w_q(\theta) = \frac{C_1^{q+1}(\theta)}{C_1^{q+1}(\theta) + C_2^{q+1}(\theta)}$. Without loss of generality, we assume $w_q(\theta_p^*) \in (0, \frac{1}{2})$. If $w_q(\theta_p^*) = \frac{1}{2}$, the gradient of norm $\mathbb{H}_{\text{norm}}$ is defined as $-\ln(\frac{w_q(\theta_p^*)}{1-w_q(\theta_p^*)}) \cdot \frac{\partial w_q(\theta_p^*)}{\partial p}$, which trivially equals to zero for any q and p . If $w_q(\theta) \in (\frac{1}{2}, 1)$, we can flip the label of C_1 and C_2 to make sure $w_q(\theta) \in (0, \frac{1}{2})$

Given $M = 2$, by applying the chain rule, the gradient of the norm can be rewritten as

$$\begin{aligned} & \frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \\ &= -\ln\left(\frac{w_q(\theta_p^*)}{1-w_q(\theta_p^*)}\right) \cdot \frac{\partial w_q(\theta_p^*)}{\partial p} \\ &= -\ln\left(\frac{w_q(\theta_p^*)}{1-w_q(\theta_p^*)}\right) \cdot \frac{\partial w_q(\theta_p^*)}{\partial \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right)^{q+1}} \cdot \frac{\partial}{\partial p} \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right)^{q+1} \\ &= -\ln\left(\frac{w_q(\theta_p^*)}{1-w_q(\theta_p^*)}\right) \cdot \left(\frac{C_2^{q+1}(\theta_p^*)}{C_1^{q+1}(\theta_p^*) + C_2^{q+1}(\theta_p^*)}\right)^2 \frac{\partial}{\partial p} \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right)^{q+1} \\ &= -\ln\left(\frac{w_q(\theta_p^*)}{1-w_q(\theta_p^*)}\right) \cdot \left(\frac{C_2^{q+1}(\theta_p^*)}{C_1^{q+1}(\theta_p^*) + C_2^{q+1}(\theta_p^*)}\right)^2 (q+1) \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right)^q \cdot \frac{\partial}{\partial p} \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right) \\ &= -\ln\left(\frac{w_q(\theta_p^*)}{1-w_q(\theta_p^*)}\right) \cdot (1-w_q(\theta_p^*))^2 (q+1) \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right)^q \cdot \frac{\partial}{\partial p} \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right). \end{aligned} \quad (27)$$

For any $q \geq 0$, it's obvious that

$$\ln\left(\frac{1-w_q(\theta_p^*)}{w_q(\theta_p^*)}\right) \cdot (1-w_q(\theta_p^*))^2 (q+1) \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right)^q \cdot \frac{\partial}{\partial p} \geq 0. \quad (28)$$

According to Eq. (17), in the point $q = p$, we have $\frac{\partial \mathbb{H}_{\text{norm}}(C^{q+1}(\theta_p^*))}{\partial p} \geq 0$, which is equivalent to

$$\frac{\partial}{\partial p} \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right) \geq 0. \quad (29)$$

Since we assume $w_q(\theta) \in (0, \frac{1}{2})$, then $C_1(\theta) < C_2(\theta)$. For any $q' \geq 0$, we also have $w_{q'}(\theta) \in (0, \frac{1}{2})$ and the following

$$\ln\left(\frac{1-w_{q'}(\theta_p^*)}{w_{q'}(\theta_p^*)}\right) \cdot (1-w_{q'}(\theta_p^*))^2 \cdot (q'+1) \left(\frac{C_1(\theta_p^*)}{C_2(\theta_p^*)}\right)^{q'} \geq 0. \quad (30)$$

By multiplying Eq. (29) with Eq. (30), for $M = 2$, we conclude for any $p \geq 0$ and $q \geq 0$,

$$\frac{\partial \mathbb{H}_{\text{norm}}(C^q(\theta_p^*))}{\partial p} \geq 0. \quad (31)$$

□

A.2.3. PROOF OF PROPOSITION 4.5

Proof. We start with a specific κ . Similar to the proof in Mohri et al. (2019), for any $\delta > 0$, the following inequality holds with probability at least $1 - \delta$ for $h \in H$:

$$\mathcal{J}_\kappa(h) \leq \mathcal{L}_\kappa(h) + \mathbb{E} \left[\max_{h \in H} \mathcal{J}_\kappa(h) - \mathcal{L}_\kappa(h) \right] + B \sqrt{\sum_m \frac{\kappa_m^2}{2N_m} \log \frac{1}{\delta}}. \quad (32)$$

Using the Hölder's inequity, we have

$$\mathcal{L}_\kappa(h) = \sum_m \kappa_m C_m \leq \left(\sum_m \kappa_m^p \right)^{\frac{1}{p}} \left(\sum_m C_m^{q+1} \right)^{\frac{1}{q+1}} = \|\kappa\|_p \tilde{\mathcal{L}}^q(h), \frac{1}{p} + \frac{1}{q+1} = 1. \quad (33)$$

Plugging $\mathcal{L}_\kappa(h) \leq \|\kappa\|_p \tilde{\mathcal{L}}^q(h)$ into Eq. (32), we obtain for $h \in H$,

$$\mathcal{J}_\kappa(h) \leq \|\kappa\|_p \tilde{\mathcal{L}}^q(h) + \mathbb{E} \left[\max_{h \in H} \mathcal{J}_\kappa(h) - \mathcal{L}_\kappa(h) \right] + B \sqrt{\sum_m \frac{\kappa_m^2}{2N_m} \log \frac{1}{\delta}}, \quad (34)$$

where $\frac{1}{p} + \frac{1}{q+1} = 1$.

Therefore, Eq. (11) in Proposition 4.5 can be readily derived from Eq. (34) by considering the maximum value across all potential κ values within Δ . $\square$

Discussions Deriving the optimal value of q that results in the tightest generalization bound from Proposition 4.5 is not trivial. In practice, our proposed Equitable Objective allows us to fine-tune a range of q values to strike a balance between performance equity/uniformity and accuracy.

A.3. Additional Experiments Details and Results

A.3.1. ADDITIONAL EMPIRICAL DETAILS

For the data centers application in Section 5.1, within each agent, the dataset is randomly partitioned, with 67% allocated as the training set and the remaining portion as the testing set. As for the EV charging application in Section 5.2, the ratio between training and testing in each agent is 70% vs. 30%. We set the learning rate as 0.05 for the data centers application and $1e-4$ for the EV charging application. We employ the Adam optimizer with a scheduler featuring a step size of 50 and a decay factor of 0.5. In both applications, the batch size is set as 128. For predicting the next time step in the data centers application, a sequence length of 12 is utilized, while in the EV charging application, the prediction involves the next charging time window spanning 12 time steps, a sequence length of 12 is also employed. The LSTM model employed in data centers application has a hidden size of 50. In the EV charging application, the Transformer model consists of a single-layer encoder-decoder with positional encoding, utilizing a feature size of 250.

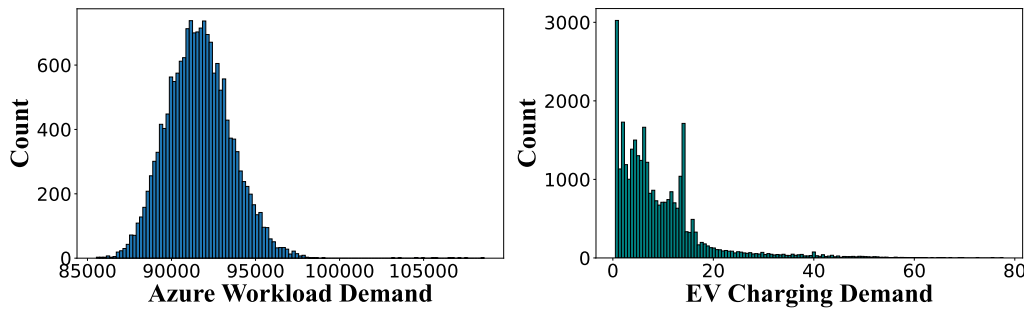


Figure 4: Depictions of (a) Azure workload demands (Shahrad et al., 2020); (b) EV charging demands in ACN-Data (Lee et al., 2019).

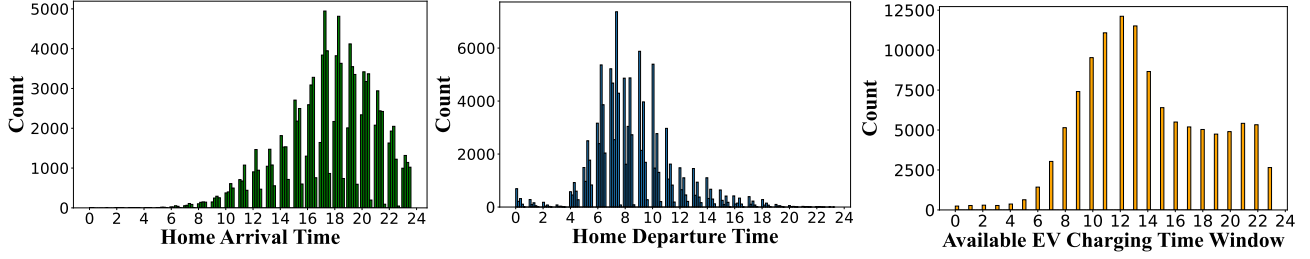


Figure 5: Depictions of home arrival, home departure and available charging time window for residential EV based on the NHTS government data (U.S. Department of Transportation, 2017).

In the EV charging scheduling application of Section 5.2, we use the publicly available National Household Travel Survey (NHTS) data (U.S. Department of Transportation, 2017) to approximate the available charging time window, *i.e.*, from s_j to e_j , for residential sectors (Wang & Paranjape, 2015). The NHTS contains the travel logs of 117, 222 American households’ vehicles, detailing the number of trips for each household and the start and end times for each trip per day. We assume the distribution for the initial charging time s_j and end time e_j of EV are the same as the distribution of home arrival and home departure times, respectively. We use the time when the last trip of a household concludes from NHTS as the daily home arrival time. Similarly, we designate the time when the first daily trip begins from NHTS as the daily home departure time (Wang & Paranjape, 2015).

A.3.2. ADDITIONAL DETAILS OF DATASETS

We depict the distribution of Azure workload demand (Shahrad et al., 2020) and EV charging electricity demands (Lee et al., 2019) in Figure 4. Additionally, in Figure 5, we present the distributions of home arrival time, home departure time, and the available EV charging time window, calculated as the difference between home departure time and home arrival time, utilizing data from the NHTS government dataset (U.S. Department of Transportation, 2017). From Figure 5, it is evident that a significant portion of residential households has an available charging time window exceeding 8 hours, thereby supporting the feasibility of scheduling environmentally friendly and financially efficient charging.

In data preprocessing of the EV charging application, we focus on the state of California to ensure alignment between the ACN-Data and CAISO. Besides null value, we also filter out the data points containing charging duration exceeding one day, as most EV can complete full charging within 5 hours, as reported by the government survey (Smart & Salisbury, 2015).

A.3.3. ADDITIONAL RESULTS

Table 2: MSE loss of the Plain PM and the EQUITABLE PM with varied $q + 1$.

	$q + 1$	MSE loss	
		Similar Agents	Different Agents
EQUITABLE PM	20	6.63	6.52
	30	6.57	6.47
	40	6.55	6.48
Plain PM	-	6.54	6.45

In completing the results of the scheduled EV charging application in Section 5.2, we report the MSE loss of each method under conditions where the distributions w.r.t. $(D_j - I_j)$ of downstream agents are “similar” and “different” in Table 2.

A.4. Experiments on Diverse Cost Objectives of Downstream Agents

We add an experiment where agents have different objective functions: *Agent (A)* for data center workload scheduling, *Agent (B)* for EV charging, and *Agent (C)* for iPhone green charging. This setup creates a diverse pool of agents with varying objectives, all utilizing carbon emission predictions from the upstream public model. The objectives for Agent (A) and (B) are defined by Eq. (12) and Eq. (13) in the main text, respectively. Note that for the EV charging application in this experiment, the public model only predicts carbon emissions E_t^C rather than E_t . For iPhone green charging, the objective is to minimize carbon emissions by optimizing the charging schedule, formulated as $\min_{X_o} \sum_t \mu_{o,t} x_{o,t} \cdot E_t^C$, where $\mu_{o,t}$

represents the electricity charged for the o -th iPhone at time t , $x_{o,t}$ is a binary variable ($x_{o,t} \in \{0, 1\}$) indicating whether charging occurs at time t , and $X_o = [x_{o,1}, \dots, x_{o,T}]$, denoting the charging schedule for the o -th iPhone.

In the implementation, we set λ to 2 for the objective of Agent (A) as indicated by Eq. (12). The dataset is split into training and testing sets with a ratio of 67% to 33%. We set the initial learning rate to 0.05 for training the Plain Public Models, and 0.1 for training EQUITABLE PM, with a step size of 50 and a decay rate of 0.1. The batch size is set to 128 for training both models. In this experiment, we use the transformer with the same architecture described in Section 5.2. For the three downstream agents with diverse objectives, we set the sequence length to 12 when predicting the next time steps of carbon emissions. In the cases of EV charging and iPhone green charging, the length of the available time frame is set to 12. For the data center application in Agent (A), which only requires the immediate next time step of carbon emission prediction, we average the predicted next 12 time steps of carbon emissions from the upstream public model. For Agent (B) and Agent (C), which need predictions for the next 12 time steps of carbon emissions, we use the predicted values directly.

We present the results of using different objectives across downstream agents in Table 3. The results indicate that even when downstream agents have distinct objective functions, our proposed EQUITABLE PM still reduces the variance in their performance distribution. This leads to a fairer solution compared to the Plain PM, which only minimizes carbon prediction error without considering the decision-making costs of diverse downstream agents.

Table 3: Statistics of the test results using different cost objectives for downstream agents.

Method	$q + 1$	Variance	Mean	$C_{95} - C_5$	MSE
EQUITABLE PM	1	18.14	7.06	9.28	9.66
	1.1	15.72	6.14	8.39	9.67
Plain PM	-	18.89	7.24	9.45	7.20

A.5. Combined Objective: Explicitly Incorporating $\mathcal{L}_f$

We present a combined objective here to complement the Equitable Objective proposed in the main text to provide a more nuanced control over equity/fairness versus model accuracy. The combined objective shown in Eq. (35) incorporates the loss of public model $\mathcal{J}_f$ into the original Equitable Objective,

$$\begin{aligned}
 & \min_{\theta} (1 - \beta) \mathcal{J}_{EQ}^q + \beta \mathcal{J}_f, \quad \text{with} \\
 & \mathcal{J}_{EQ}^q = \sum_{m=1}^M \mathbb{E}^{q+1} [\text{cost}_m(\hat{a}_m, \xi_m, y) - \text{cost}_m(a_m, \xi_m, y)] \\
 & \mathcal{J}_f = \mathbb{E}[\|y - \hat{y}\|^2],
 \end{aligned} \tag{35}$$

where β controls the weighting of each component. We then approximate the expectation in Eq. (35) with the empirical loss as shown in the Eq. (36).

$$\begin{aligned}
 & \min_{\theta} (1 - \beta) \mathcal{L}_{EQ}^q + \beta \mathcal{L}_f, \quad \text{with} \\
 & \mathcal{L}_{EQ}^q = \sum_{m=1}^M \left\{ \left[\frac{1}{N_m} \sum_{i=1}^{N_m} (\text{cost}_m(\hat{a}_{m,i}, \xi_{m,i}, y_{m,i}) - \text{cost}_m(a_{m,i}, \xi_{m,i}, y_{m,i})) \right]^{q+1} \right\} \\
 & \mathcal{L}_f = \sum_{m=1}^M \frac{1}{N_m} \sum_{i=1}^{N_m} \|y_{m,i} - \hat{y}_{m,i}\|^2
 \end{aligned} \tag{36}$$

Likewise, if the cost functions are differentiable, the gradient of the combined objective is calculated as

$$\begin{aligned}
 \nabla_{\theta}((1 - \beta) \mathcal{L}_{EQ}^q + \beta \mathcal{L}_f) = & (1 - \beta) \sum_{m=1}^M \sum_{i=1}^{N_m} \nabla_{C_{m,i}} \mathcal{L}_{EQ}^q \nabla_{\hat{a}_{m,i}} C_{m,i} \nabla_{\hat{y}_i} \hat{a}_{m,i} \nabla_{\theta} \hat{y}_{m,i} \\
 & + \beta \sum_{m=1}^M \sum_{i=1}^{N_m} \frac{2}{N_m} (\hat{y}_{m,i} - y_{m,i}) \nabla_{\theta} \hat{y}_{m,i},
 \end{aligned}$$

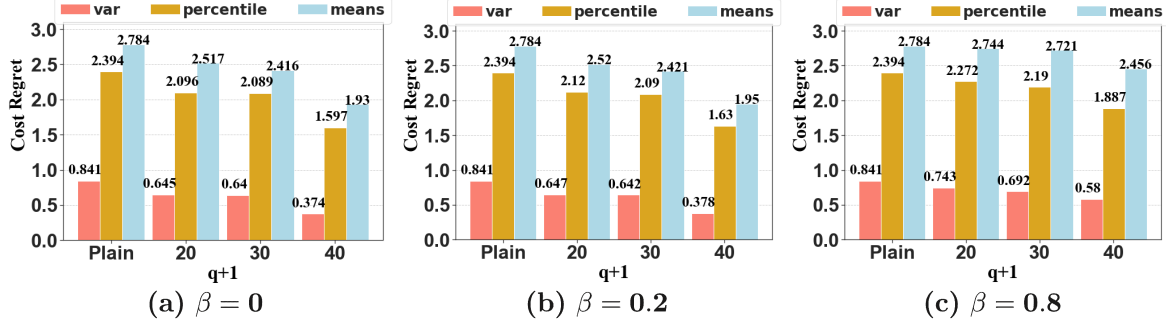


Figure 6: Statistics of test results when $\beta = [0, 0.2, 0.8]$. Note the EQUITABLE PM here refers to the combined objective, Eq. (35). We can observe that the EQUITABLE PM has achieved more uniform distributions among agents compared to the Plain PM, according to the variance and percentile difference measures.

where $C_{m,i} = \text{cost}_m(\hat{a}_{m,i}, \xi_{m,i}, y_{m,i}) - \text{cost}_m(a_{m,i}, \xi_{m,i}, y_{m,i})$.

If the cost functions are non-differentiable, similar as (4), given a training dataset with K batches and a batch size B_m , the gradient can be calculated as

$$\nabla_{\theta} \mathcal{L}^q(\theta) = \frac{1}{K} \sum_{k=1}^K \left\{ \left[\sum_{i=1}^{B_m} \sum_{m=1}^M \nabla_{\theta} \log \sigma_{\theta}(\hat{y}_{m,k,i} | x_{m,k,i}) \right] \cdot \left[\sum_{m=1}^M \left[(1 - \beta) \left(\frac{1}{B_m} \sum_{i=1}^{B_m} C_{m,k,i} \right)^{q+1} + \beta \sum_{i=1}^{B_m} \frac{1}{B_m} \mathcal{L}_{f,m,i} \right] \right] \right\}. \quad (37)$$

It is not straightforward to prove that a larger q would lead to a more uniform cost regret distribution by the combined objective. The challenge arises because θ' that minimizes $\mathcal{L}_{EQ}^q$ may not align with θ^* that optimizes the combined loss of $\mathcal{L}_{EQ}^q$ and $\mathcal{L}_f$. Nevertheless, we highlight that the combined objective provides a way to allow us to balance between fairness of downstream agents and upstream public model accuracy, achieved by adjusting the value of β .

Table 4: MSE loss of the Plain PM and the EQUITABLE PM with $q + 1$ as 40. Note the EQUITABLE PM here refers to the combined objective, Eq. (35). As β increases, the MSE loss of EQUITABLE PM decreases.

MSE loss		
EQUITABLE PM	$\beta = 0$	6.48
	$\beta = 0.2$	6.48
	$\beta = 0.8$	6.46
Plain	-	6.45

A.5.1. EMPIRICAL RESULTS FOR THE COMBINED OBJECTIVE

We perform empirical investigations under the same setup outlined in Section 5.2 to examine whether the proposed combined objective in Eq. (36) could lead to a more equitable performance distribution among agents in the EV Charging Scheduling case study. Various β and q values are considered. Note the EQUITABLE PM mentioned in the following results refers to the public model trained using the combined objective in Eq. (36).

Results Figure 6 reports the evaluation results between the Plain PM and EQUITABLE PM with different q and β values. It can be observed the variance and $C_{95} - C_5$ achieved by the EQUITABLE PM consistently remain lower than using the Plain PM. From Figure 6, we observe both the variance and $C_{95} - C_5$ of cost regret distributions across agents decreases as the value of q increases, implying the performance distribution becomes more uniform. Notably, setting $\beta = 0$ makes the EQUITABLE PM focus on optimizing the Equitable Objective $\mathcal{L}_{EQ}^q$ exclusively, resulting in the most uniform distribution compared to $\beta = 0.2$ and $\beta = 0.8$. In contrast, the MSE loss, $\mathcal{L}_f$, decreases as β increases, as shown in Table 4.