
PreferenceNet: Encoding Human Preferences in Auction Design with Deep Learning

Neehar Peri*, Michael J. Curry*, Samuel Dooley, John P. Dickerson

Center for Machine Learning, University of Maryland

peri@umiacs.umd.edu, {curry, sdooley1, john}@cs.umd.edu

Abstract

The design of optimal auctions is a problem of interest in economics, game theory and computer science. Despite decades of effort, strategyproof, revenue-maximizing auction designs are still not known outside of restricted settings. However, recent methods using deep learning have shown some success in approximating optimal auctions, recovering several known solutions and outperforming strong baselines when optimal auctions are not known. In addition to maximizing revenue, auction mechanisms may also seek to encourage socially desirable constraints such as allocation fairness or diversity. However, these philosophical notions neither have standardization nor do they have widely accepted formal definitions. In this paper, we propose PreferenceNet, an extension of existing neural-network-based auction mechanisms to encode constraints using (potentially human-provided) exemplars of desirable allocations. In addition, we introduce a new metric to evaluate an auction allocations’ adherence to such socially desirable constraints and demonstrate that our proposed method is competitive with current state-of-the-art neural-network based auction designs. We validate our approach through human subject research and show that we are able to effectively capture real human preferences. Our code is available on GitHub

1 Introduction

Auctions are an essential tool in many marketplaces, including those in electricity [8], advertising [14], and telecommunications [9, 26]. The design of auctions with desirable properties is thus an important theoretical and practical problem. A typical assumption is that bidders may choose to bid strategically and will successfully anticipate the behavior of other bidders. This results in a potentially complicated Bayes-Nash equilibrium which may be difficult to predict. To evade this problem, a common requirement is that an auction should be strategyproof (i.e. bidders should be incentivized to truthfully share their valuations regardless of other bidders’ actions).

If the goal of an auction is to maximize the total welfare of all participants, the Vickrey-Clarke-Groves (VCG) mechanism is both strategyproof and welfare-maximizing [40, 7, 20]. Intuitively, in many cases the VCG mechanism corresponds to a second-price auction. If the goal is to maximize the revenue gained, the problem is more challenging. Myerson’s work completely characterizes strategyproof, revenue-maximizing auctions for a single item [29]. Beyond this case, results are more limited. Some results are known for the “multiple-good monopolist” problem, in which the auctioneer sells multiple items to a single bidder [11, 12, 24, 31, 28, 21, 18]. In addition, designing auctions for the related but weaker notion of Bayes-Nash incentive compatibility is reasonably well understood [2, 3, 4].

There has been significant difficulty in designing strategyproof auctions involving multiple items and multiple bidders. [41] shares significant, but limited results which solve the case in which items

* The first two authors contributed equally to this work. 35th Conference on Neural Information Processing Systems (NeurIPS 2021), Sydney, Australia.

may have at most 2 values. Due to the apparent difficulty of analytically designing strategyproof, revenue maximizing auctions, recent methods instead approximate optimal auctions using machine learning approaches [13, 10, 19, 25, 15, 33, 32]. [37] proposes an method which guarantees exact strategyproofness in the single-agent setting. [39, 1] also use neural networks in the design of auctions. These methods primarily focus on standard mechanism design goals of revenue or welfare maximization, but these may not be the only goals of an auction. An auction might be conducted for public goods (i.e. electromagnetic spectrum distribution [9]), where the auctioneer must consider the effect not only on the auctioneer and bidders but on third-parties as well [30]. In other cases, such as auctions involving job or credit advertisements, it might be necessary to additionally constrain the auction mechanism to ensure fairness with respect to protected characteristics [23, 6]. Recent work has considered the problem of determining revenue-maximizing, strategyproof, fair auctions, either from a specific class [5] or with a general neural network approach [25].

The underlying notions of “fairness” used in these papers (e.g. total-variation fairness [23]) are defensible but somewhat arbitrary – they are attempts to formally capture human intuition. It might instead be better to attempt to capture, from data, human intuitions about fairness or other ethical requirements that are not explicitly defined [38, 16].

In this paper, we introduce PreferenceNet, an extension to RegretNet that directly encodes socially desirable constraints from data, and captures noisy human preferences. We are motivated by advertising auctions where the allocations of the auction mechanism must satisfy human preference constraints alongside the typical goals of strategyproofness and revenue maximization. We conduct a number of experiments using synthetic data (as is typical for neural network based auction mechanisms [13]) and evaluate our method on different auction settings and fairness constraints. We show that PreferenceNet is able to effectively capture each fairness constraint and matches the performance of standard approaches. We also conduct two surveys to further study human preferences. In the first survey (n=140), when given a specific definition of fairness, we ask participants to determine if a given auction setting is fair in order to test the noise in human judgements. In the second survey (n=345), we elicit judgements of pairwise comparisons of two auction settings to determine preferences without priming the participants with any particular definition of fairness. We train PreferenceNet on these data and show that our approach can capture nuanced human preferences in auction design.

2 Background and Related Work

RegretNet. [13] presents RegretNet, a neural-network architecture for learning approximately strategyproof auctions that maximize revenue. RegretNet treats the auction mechanism as a function from bids to allocations and payments, parameterized as a neural network. Revenue is optimized via gradient descent on sampled truthful bid profiles; strategyproofness is enforced by computing strategic bids in an adversarial manner to minimize violations. RegretNet has been modified and extended in a variety of ways, particularly to enforce additional desirable constraints. [10, 39, 19, 25, 15, 33].

Special Case: Single-Bidder Auctions. We highlight a special case in which deep learning for auctions has been particularly successful – auctions with a single bidder. In the single-bidder setting, the set of strategyproof mechanisms can be easily characterized [34], and it is possible to design neural network architectures which will always lie in this set. [13, 37] present two different learning-based solutions. Both methods are guaranteed to be strategyproof, and revenue can be maximized by unconstrained optimization. There are some known optimal single-agent mechanisms (we highlight those of Manelli-Vincent [28] and Pavlov [31], as they are relevant to auction settings we study below). Moreover, the theory of single-bidder mechanisms is well understood [11, 12, 18], so it is also possible to learn empirically strong mechanisms using these approaches and then prove their optimality [37, 13]. Unfortunately, when moving beyond the single-agent setting, it is necessary to use more general neural network architectures for which these guarantees do not apply. Because we are interested in such settings, we focus on these general architectures in this work.

Fairness and Human Preferences. As mentioned, while revenue, strategyproofness, and individual rationality are the classic goals of auction design, it might also be necessary for allocations made by an auction to satisfy certain other requirements such as fairness [5, 25]. However, it is not always

clear if these mathematical definitions of these concepts actually capture human intuitions. [36, 38] considers human reactions to different definitions of fairness in a classification setting. [35] tests the extent to which human participants are able to understand and apply fairness metrics. [16] considers the problem of learning to perform fair clustering from human-provided demonstrations. As discussed below, we also crowdsource human opinions on fairness in auctions and analyze the results in Section 6.

3 Problem Setting

Auction Model. An auction is defined as a set of agents $N = \{1, \dots, n\}$ bidding for items $M = \{1, \dots, m\}$. Each agent $i \in N$ has a corresponding private valuation v_i , randomly drawn from a set of n valuations as $v = (v_1, \dots, v_n) \in V_i$. In general v_i may be functions over the power set of items 2^M . However, we only consider simpler cases with **additive** valuations and **unit-demand** valuations, where the valuation is simply a vector $v_i \in \mathbb{R}^m$ of values per item.

Each agent reports a bid vector b_i to the auctioneer, which may differ from the private valuation v_i . Given the profile of bids $b = (b_1, \dots, b_n)$ of all agents, the auction has allocation and payment rules $g(b) : \mathbb{R}^{mn} \rightarrow [0, 1]^{nm}$ and $p(b) : \mathbb{R}^{mn} \rightarrow \mathbb{R}^n$. We will refer to the matrix of allocation probabilities, whose rows must sum to 1, as $g(b) = z$. Likewise agent i 's value of the j th item is $v_{i,j}$, and bidder i has payment function p_i . Moreover, for unit-demand auctions, we restrict the allocation to allow each bidder to win, in expectation, at most 1 item. Given the allocation, each agent receives a utility which can in either case be represented in linear form as $u_i(v) = \sum_j v_{i,j} z_{i,j} - p_i$.

Desirable Auction Properties. A mechanism is individually rational (IR) when an agent is guaranteed non-negative utility: $u_i(v_i; v) \geq 0 \forall i \in N, v \in V$. A mechanism is dominant-strategy incentive-compatible (DSIC) or strategyproof if every agent maximizes their own utility by bidding truthfully, regardless of the other agents' bids. We can define regret, the difference in utility between the bid player i actually made and the best possible strategic bid: $\text{rgt}_i(v) = \max_{b_i} u_i(b_i, v_{-i}) - u_i(v_i, v_{-i})$. In addition to satisfying the IR and DSIC constraints, the auctioneer seeks to maximize their expected revenue. If the auction is truly DSIC, players will bid truthfully, and as a result revenue is simply $E_{v \sim V} [\sum_{i \in N} p_i(v)]$.

4 PreferenceNet: Encoding Human Preferences

We first explore a new metric to evaluate the adherence to socially desirable constraints in item allocations. Next, we describe the implementation details of PreferenceNet and important considerations when training the model.

Evaluation Metrics. There are limited evaluation criteria that quantitatively measures an auction mechanism's ability to enforce constraints on item allocations. If the constraints are not known explicitly, one can qualitatively examine the allocation graphs to evaluate the underlying allocation function $g(b)$. However, visual inspection does not scale to larger auction settings. As shown in Figure 1, our proposed metric is not only able to capture the same insights as qualitative analysis, but also scales to arbitrarily large auction mechanisms.

Given the limitations of existing analysis techniques, we propose Preference Classification Accuracy (PCA), a new metric to evaluate how well a learned auction model satisfies an arbitrary constraint. For a set of test bids b and allocations $g(b) : \mathbb{R}^{mn} \rightarrow [0, 1]^{nm}$ generated by our learned auction model, we assign a label $s(b) \in \{1, 0\}$ to each allocation according to a ground truth labeling function based on the underlying preference. For each test bid b , $s(b)$ is 1 if the learned auction network satisfies the ground truth constraint. PCA is calculated by averaging the value of $s(b)$ over n test bids. Note that this metric remains valid in cases when we know the underlying preference function (e.g. total variation fairness, entropy) as well as when we are sampling from an unknown distribution (e.g. human preference elicitation). For cases where the underlying preference function is known, we can directly compute the preference score for a given allocation and apply a threshold to obtain a label. For cases where the preference function is unknown, as is the case in human preference elicitation, we can use the ground truth allocation-label pair to assign preference labels $s(b)$ to new allocations based on the nearest neighbor in the ground truth set. This metric gives us a formal way to measure

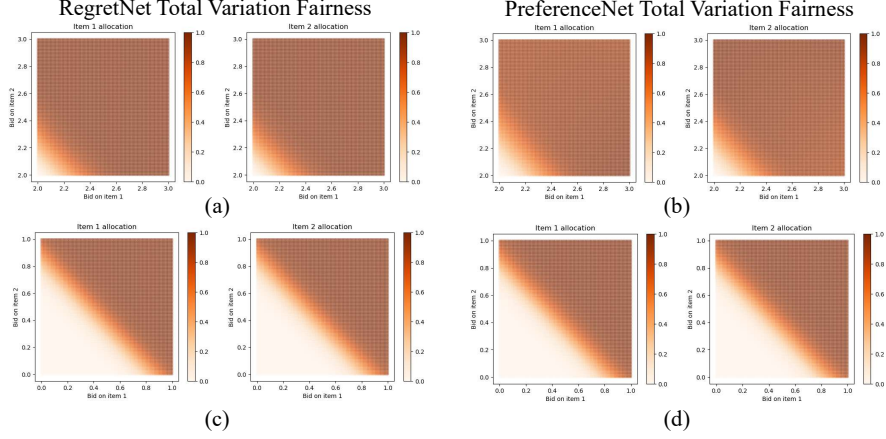


Figure 1: We compare the allocation plots of standard RegretNet approaches in enforcing total variation fairness (TVF) [25, 23] (sub-figures a and c) with our proposed approach (sub-figures b and d) learned through exemplars of desirable allocations that satisfy TVF. Visually, the allocations for both the unit-demand (sub-figures a and b) and additive auctions (sub-figures c and d) are identical. In this case, our proposed metric verifies our visual inspection, indicating that allocations from all four models satisfy TVF with 100% accuracy. However, this qualitative analysis does not extend to larger auction settings. In contrast, our proposed metric allows us to quantify the adherence of an auction mechanism to an enforced constraint for arbitrarily large auctions.

the degree to which preference constraints are violated, which crucially can be used whether or not the constraints follow an explicitly-known function.

Preference Elicitation. In order to effectively elicit preferences, we rely on pairwise comparisons between allocations to identify both positive and negative exemplars. We compare each input set of allocations against n other allocations to determine if a particular sample is preferred over these alternatives (see Figure 2). This method of group preference elicitation reduces noise and ensures that the learned preference is satisfactory to a majority of the participants. We use this ranking approach to generate training labels in all of our experiments as described below.

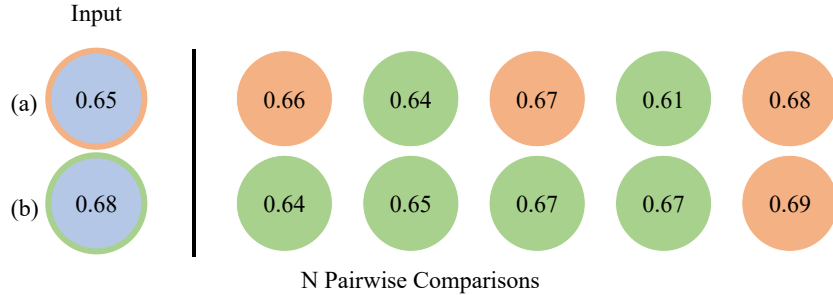


Figure 2: To elicit preferences from a group, or from a person without deterministic preferences, we use pairwise comparisons to determine labels for each allocation. In examples (a) and (b), each allocation (represented as a circle) has a preference score (represented by a number inside the circle). We compare the score of the input against n other valuations to determine the relative ranking of the input data point. If the input data point has a smaller score than a plurality of the points it is compared against, it is a negative exemplar for the implicit preference (as in (a)). Otherwise it is positive (as in (b)).

Training Algorithm. We present the training algorithm of PreferenceNet below. We use the same additive and unit-demand network architectures as RegretNet for arbitrary numbers of agents and items. Our training algorithm follows closely from RegretNet. PreferenceNet consists of two sub-networks: RegretNet and a 3-layer MLP with a sigmoid activation at the output. We first train the MLP using a uniformly drawn sample of allocations as inputs, and optimize the binary cross entropy loss function using ground truth labels (generated as in Figure 2) identifying positive and negative exemplars. Next, we train RegretNet using the standard training procedure with the modified loss function described in Eq. 1. Lastly, we sample allocations and payments from the partially trained RegretNet model every c epoch and augment the MLP training set to adapt to the distributional shifts in allocations over the course of training. Our modified loss function $\mathcal{C}_\rho(w; \lambda)$ is defined as:

$$\begin{aligned} \mathcal{L}_{\text{rgt}} &= \sum_{i \in N} \lambda_{(r,i)} \text{rgt}_i(w) + \frac{\rho_r}{2} \sum_{i \in N} \text{rgt}_i(w)^2, \quad \mathcal{L}_s = \sum_{j \in M} s_j \\ \mathcal{C}_\rho(w; \lambda) &= -\frac{1}{L} \sum_{l=1}^L \sum_{i \in N} p_i^w(v^{(l)}) + \mathcal{L}_{\text{rgt}} - \mathcal{L}_s \end{aligned} \quad (1)$$

where $\mathcal{L}_s$ is the output of the trained MLP. RegretNet is optimized such that it is strategyproof, revenue-maximizing, satisfies the preference learned by the MLP (i.e. maximizes the output scores of the MLP).

For each configuration of n agents and m items, we train RegretNet for a maximum of 200 epoch using 160,000 training samples. We also train the MLP with 80,000 initial training samples and iteratively retrain the MLP with 5,000 additional samples from the partially trained RegretNet every 5 epoch. For all networks, we use the Adam optimizer and 100 hidden nodes per layer. We apply warm restarts to the MLP optimizer each time we add new training data to prevent the model from settling in a local minima. We incremented ρ_r every 2500 iterations and λ_r every 25 iterations. Finally, we report the preference classification accuracy, mean regret, and mean payments by simulating the allocations of 20,000 testing samples. We run all our experiments on an NVIDIA Titan X (Pascal) GPU. We refer readers to GitHub for our implementation.

MLP Architecture. We learn implicit preference functions using a simple 3-layer multi-layer perceptron that takes as input a mini-batch of allocations and outputs a vector that scores the input $\in [0, 1]$ as a measure of how closely the allocation satisfies the ground truth preference. Given that neural networks are universal function approximators, with enough parameters the MLP can represent any arbitrary preference function. In practice, we find that ReLU non-linear activations and batch normalization are essential to effectively train this network.

Class Balanced Sampling. The distribution of positive and negative training examples for arbitrary preferences are often imbalanced – preferred examples may be concentrated in a small region of possible allocations. Often, this imbalance can make it difficult to train a robust model. Commonly, neural networks with improper class balance fail to learn a good decision boundary. In order to compensate for such imbalanced datasets, we can explicitly over-sample sparse classes until there are an equal number of positive and negative training examples. In practice this allows the network to quickly learn the decision boundary and allows us to train with fewer data points.

MLP Co-Training. The distribution of payments and allocations generated by RegretNet shift significantly during the course of training. As a result, the initially trained MLP may not effectively enforce the preference loss as RegretNet continues to train. In order to adapt to the distribution shift, we sample allocations from RegretNet at fixed intervals while it is training, add them back to the training set, and retrain the MLP. We generate noisy labels [27] using the existing MLP to reduce the cost of collecting expensive labels. In practice, we find that this approach effectively reinforces the decision boundary between positive and negative exemplars.

Model Validation and Selection. The optimal model minimizes regret, while maximizing payments and preference classification accuracy. In practice, satisfying all three conditions may be difficult. Often, we find that choosing a fixed epoch to evaluate does not provide consistent results. Rather, we evaluate each checkpoint on a validation set and maximize the following criteria in Eq. 2:

$$\alpha * \text{PCA} + \beta * \frac{p(\bar{b})}{\max(p(b))} + \gamma * \left(1 - \frac{\text{rgt}(\bar{b})}{\max(\text{rgt}(b))}\right) \text{ s.t. } \alpha + \beta + \gamma = 1 \quad (2)$$

In our experiments we set $\alpha = 0.45, \beta = 0.1, \gamma = 0.45$. It is important to note that the maximum regret and payments are calculated over the entire training process, while the mean regret, payment, and preference classification accuracy are calculated at each epoch.

5 Sampling Synthetic Preferences

Given the lack of publicly available auction data, we generate synthetic bids as in [13, 25, 37]. We first validate PreferenceNet using synthetic preferences, and extend our analysis to real human preferences in Section 6. In our synthetic preference experiments, we are interested in three types of fairness: total variation fairness (TVF), entropy, and a quota system. All three definitions of fairness map the allocations $g(b)$ onto $\mathbb{R}$. We train auction models for each of these valuation function and compare RegretNet with our proposed model under uniform additive and unit-demand auction settings in Table 5. Note in Table 5 that RegretNet is trained with an additional loss function term to explicitly optimize for the preference. PreferenceNet is trained as described in 4. We describe the three definitions of fairness below:

Total Variation Fairness. An auction mechanism satisfies total variation fairness if the ℓ_1 -distance between allocations for any two users is at most the distance between those users. That is, total variation fairness is satisfied when

$$\forall k \in \{1, \dots, c\}, \forall j, j' \in M, \sum_{i \in C_k} |g(b)_{i,j} - g(b)_{i,j'}| \leq d^k(j, j'). \quad (3)$$

We fix $d = 0$ in all of our experiments. We minimize violations of the above constraints.

Entropy. An auction mechanism satisfies the entropy constraint if the entropy of the normalized allocation for a bid profile b

$$-\sum_{i=1}^n P \left(\frac{g(b)_{i,\cdot}}{\sum_j g(b)_{ij}} \right) \log P \left(\frac{g(b)_{i,\cdot}}{\sum_j g(b)_{ij}} \right) \quad (4)$$

is maximized. We normalize across items, turning the allocation into a probability distribution. This normalization ensures that entropy reflects diversity in the allocations, and not the overall number of items being allocated. Specifically, encouraging entropy ensures that the allocation for a given agent will tend to be more uniformly distributed.

Quota. An auction mechanism satisfies the quota constraint if, for each item in the (normalized) allocation, the smallest allocation to any agent j is greater than some minimum threshold t :

$$\min_j \left(\frac{g(b)_{\cdot,j}}{\sum_i g(b)_{i,j}} \right) > t \quad (5)$$

Here, we normalize the allocation so that the allocation per item will be a probability distribution over agents. The intuition for this definition is that each agent must guarantee some floor across every item. Returning to the advertising example, this ensures some minimum percentage of ad impressions are seen by every demographic group.

We train a number of models to compare the performance between RegretNet and PreferenceNet. Despite leveraging an weaker, implicit signal to learn fairness constraints, PreferenceNet is able to closely match the performance of RegretNet. Surprisingly, PreferenceNet improves upon RegretNet in some cases, indicating that with careful hyperparameter tuning, further improvements are possible. Of the three fairness constraints examined, enforcing a quota is hardest for additive auctions, as both RegretNet and PreferenceNet struggled for auctions with a large number of agents.

Limitations. Despite its effectiveness, we highlight limitations of our approach. In general, PreferenceNet always optimizes for the simplest function. For example, if learning a piece-wise preference where the positive exemplars are not clustered in a single region as in Figure 3, PreferenceNet tends to only satisfy part of the piece-wise function. This is unsurprising, given that neural networks are known to take shortcuts in optimization [17]. Moreover, PreferenceNet has difficulty in generating allocations with tightly clustered preference scores to satisfy a particular constraint. Given that we optimize for preferences implicitly using exemplars, this behavior is understandable. We study these limitations further in the supplemental material.

Table 1: We evaluate both RegretNet and PreferenceNet over $n \times m$ auctions (“u” refers to unit-demand and “a” refers to additive), where n is the number of agents and m is the number of items. We measure three criteria (a) **PCA**, (b) **Regret Mean (STD)**, (c) **Payment Mean (STD)**. Although PreferenceNet learns each preference implicitly, it produces similar performance to our strong baseline.

(a)	TVF		Entropy		Quota	
	RegretNet	PreferenceNet	RegretNet	PreferenceNet	RegretNet	PreferenceNet
2x2 u	100.0	100.0	86.4	69.6	100.0	100.0
2x4 u	100.0	100.0	99.7	94.9	100.0	100.0
4x2 u	100.0	100.0	100.0	94.5	.1	100.0
4x4 u	100.0	99.3	100.0	67.4	100.0	100.0
2x2 a	100.0	100.0	99.6	99.5	75.1	100.0
2x4 a	100.0	100.0	100.0	100.0	36.0	100.0
4x2 a	100.0	100.0	99.9	100.0	.1	.1
4x4 a	100.0	99.9	100.0	96.1	0	0

(b)	TVF		Entropy		Quota	
	RegretNet	PreferenceNet	RegretNet	PreferenceNet	RegretNet	PreferenceNet
2x2 u	.012 (.012)	.012 (.012)	.02 (.02)	.011 (.01)	.012 (.013)	.013 (.012)
2x4 u	.008 (.006)	.016 (.013)	.045 (.045)	.022 (.019)	.012 (.01)	.034 (.022)
4x2 u	.015 (.009)	.028 (.013)	.025 (.025)	.056 (.018)	.026 (.017)	.819 (.136)
4x4 u	.033 (.024)	.067 (.037)	.031 (.031)	.029 (.014)	.037 (.023)	.432 (.145)
2x2 a	.005 (.004)	.006 (.004)	.006 (.006)	.013 (.008)	.008 (.007)	.05 (.031)
2x4 a	.008 (.011)	.008 (.009)	.007 (.007)	.008 (.009)	.01 (.01)	.078 (.032)
4x2 a	.038 (.076)	.014 (.011)	.036 (.036)	.162 (.052)	.017 (.013)	.017 (.013)
4x4 a	.424 (.324)	.139 (.104)	.015 (.015)	.257 (.1)	.038 (.017)	.039 (.018)

(c)	TVF		Entropy		Quota	
	RegretNet	PreferenceNet	RegretNet	PreferenceNet	RegretNet	PreferenceNet
2x2 u	4.18 (.45)	4.17 (.44)	4.26 (.37)	4.14 (.43)	4.19 (.35)	4.18 (.33)
2x4 u	4.37 (.29)	4.41 (.34)	4.48 (.34)	4.47 (.21)	4.44 (.16)	4.49 (.29)
4x2 u	4.93 (.21)	4.98 (.21)	4.85 (.23)	4.99 (.21)	5.16 (.24)	5.03 (.21)
4x4 u	8.81 (.39)	8.92 (.38)	8.79 (.5)	8.81 (.42)	8.8 (.3)	8.81 (.36)
2x2 a	.87 (.31)	.87 (.32)	.88 (.32)	.9 (.31)	.73 (.37)	.6 (.3)
2x4 a	1.75 (.38)	1.74 (.4)	1.77 (.45)	1.74 (.4)	1.76 (.44)	1.39 (.5)
4x2 a	1.1 (.34)	1.2 (.22)	1.1 (.34)	1.15 (.22)	1.3 (.23)	1.31 (.23)
4x4 a	2.58 (.39)	2.41 (.36)	2.26 (.3)	2.28 (.32)	2.52 (.32)	2.56 (.33)

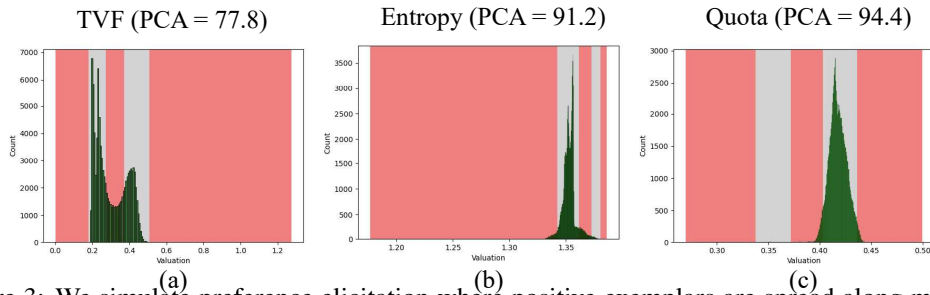


Figure 3: We simulate preference elicitation where positive exemplars are spread along multiple bands. The grey bands represent ground truth regions of positive exemplars, and the red bands represent ground truth regions of negative exemplars. In each plot, the green histogram represents the preference scores of the generated allocations. A model that perfectly enforcing a given preference rule will generate allocations that have valuations entirely within the grey bands.

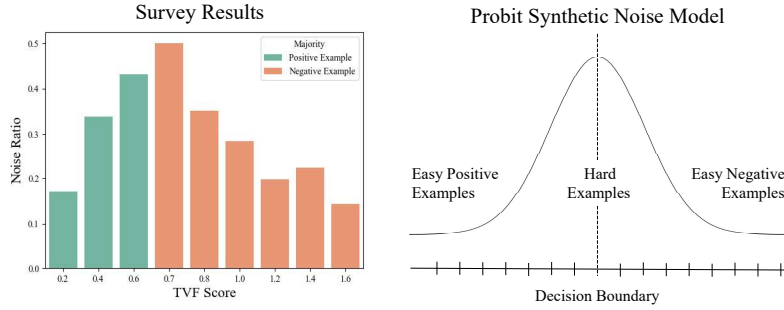


Figure 4: We crowdsource human annotators to examine various advertising scenarios and determine if an allocation is fair according to a given definition. We simplify the definition of total variation fairness (TVF) and measure the noise in responses as a function of the ambiguity of the scenario. We expect that the the label noise for a particular TVF value is inversely proportional to the distance from the decision boundary. Concretely, easy allocations will have lower noise ratios, and ambiguous allocations will have high noise ratios. We can leverage this model of uncertainty and apply it to our synthetic experiments to better simulate human preference elicitation.

6 Soliciting Human Preferences

In this section, we detail the creation of and results derived from two human subjects surveys. The results of these surveys are used to validate core assumptions of our model and provide data that we use to train an auction model. We find that PreferenceNet is able to adhere to the notion of fairness expressed by the human’s preferences of auction outcomes. This encouraging result validates the utility and expressiveness of PreferenceNet.

We conducted both surveys through Cint, a crowdsourcing platform which connected us with English-speaking participants located in the United States. After submitting an application for our human subjects research to our institutional IRB, we were notified that the survey was exempt from IRB review. Our survey protocols can be found in the supplemental materials. Cint compensates per survey completion (regardless of length to complete), and both surveys were set to pay \$1.92. All survey results are anonymized to protect participant privacy. We include these results in the supplemental material.

Measuring Preference Noise. Our first survey was designed to test how participants would interpret and apply a simple fairness definition. The survey protocol was designed as follows: We primed participants to consider an advertising auction that was shown to two different (vague) demographic groups. Each participant was told that an auction would be fair if each ad was presented to each group at equal rates. After some familiarization, we asked the participants to determine if a given scenario was fair. Each participant was asked 30 such questions. The median completion time for the survey was 6 minutes with a median hourly wage of \$18. The 30 questions presented to the participants came from a question bank of 64 randomly generated scenarios, each with an associated TVF score.

Human understanding is an inherently noisy process. Despite providing the same context, we observe that survey participants understand a given definition of fairness and apply it to various scenarios differently. Given this data, we perform a normality test using a Q-Q plot as shown in the supplemental material. We find that our survey data are well correlated with the Gaussian distribution. This is also well supported by our visual inspection of the noise distribution. Experimentally, we observe that participants’ choice of what is fair has a decision boundary at approximately a TVF value of 0.7. Interestingly, the noise is maximal near the decision boundary as shown in Figure 4. Concretely, we can model this noise using a probit model, so that the probability that the label is unperturbed for a particular TVF value increases with distance such that $P(Y = 1|X) = k\Phi(\frac{|x-\mu|}{\sigma})$, where Φ is the CDF of the Gaussian distribution, μ is the decision boundary, σ is the measured sample standard deviation, and k is an optional parameter that can scale the noise. We can use this noise model to perturb the input training data to the MLP to better simulate real data. We explore this further in the supplemental material.

Unit Demand Auction					Additive Auction				
	Human	TVF	Entropy	Quota		Human	TVF	Entropy	Quota
Human	0	0.005	0.004	0.009	Human	0	0.002	0.002	0.032
TVF	0.005	0	0.002	0.011	TVF	0.002	0	0.001	0.033
Entropy	0.004	0.002	0	0.010	Entropy	0.002	0.001	0	0.032
Quota	0.009	0.011	0.010	0	Quota	0.032	0.033	0.032	0

Human Preferences	PCA	Regret Mean (STD)	Payment Mean (STD)
2x2 Unit	100.0	.016 (.015)	4.20 (.37)
2x2 Additive	100.0	.005 (.004)	.87 (.31)

Figure 5: We randomly sample 20,000 bids and compute the average L2 distance from each model’s learned allocations to identify allocation similarity. We find that human preferences are most similar to both TVF and entropy in both the unit demand and additive auction settings. Moreover, PreferenceNet is able to perfectly capture human preferences with low regret.

Group Preference Elicitation. We designed our second survey to ask similar questions to the preference noise survey above. However, there were two primary differences: (1) we did not prime the participants with a fairness definition, and (2) the participants were presented with two scenarios and were asked which they thought was more fair. Each participant was asked 30 of these pairwise questions on the questions described above (full protocol details in the supplemental materials). This survey had 345 participants who had a median completion time of 7 minutes and a median pay of \$15 per hour.

Notions of fairness and diversity have neither standardized nor widely accepted formal definitions [38, 35, 36, 22]. The purpose of this survey is to elicit preferences from a group and train an auction model whose allocations resemble the group preference. Using the survey data, we apply our preference elicitation strategy as described in Section 4 to generate training labels for each sample. After training PreferenceNet to enforce the group preference for both the unit-demand and additive auction settings, we find that the group preference is more similar to TVF and entropy. As shown in Figure 5, human preferences are not perfectly captured by these typical models, both because group preferences rarely converge to a unifying model, and preference elicitation is a noisy process.

7 Conclusion

Although surveying people to elicit their preferences can effectively help us model ambiguous definitions of socially beneficial constraints, we must be careful about the framing of the survey questions and the choice of audiences we survey.

Ethical Implications. Humans are inherently biased, so we need to be cognizant of the effects these latent biases might have over preferences for fairness. Moreover, sampling human preferences facilitates opportunities for data poisoning attacks, in which a malicious survey respondent could try to negatively impact the survey collection process. In general, we can mitigate both of these issues by sampling at scale to avoid problems with noisy labels, although this brings additional cost. Most importantly, we must involve stakeholders to ensure that their preferences are validated through the learned model in an iterative process.

In this paper we present PreferenceNet, a novel extension to RegretNet that makes it easier to learn preferences from data to encode socially desirable constraints for auction design. We introduce a new metric to empirically measure how closely a learned mechanism enforces a particular constraint, and show that our proposed method is able to effectively capture human preferences.

Acknowledgments and Disclosure of Funding

This research was supported in part by NSF CAREER Award IIS-1846237, NSF D-ISN Award #2039862, NSF Award CCF-1852352, NIH R01 Award NLM-013039-01, NIST MSE Award #20126334, DARPA GARD #HR00112020007, DoD WHS Award #HQ003420F0035, and a Google Faculty Research Award. We thank the authors of ProportionNet for sharing their codebase and Kevin Kuo and Uro Lyi for their feedback in writing this paper.

References

- [1] Gianluca Brero, Benjamin Lubin, and Sven Seuken. “Machine Learning-powered Iterative Combinatorial Auctions”. In: (Nov. 2019). arXiv: 1911.08042 [cs.GT].
- [2] Yang Cai, Constantinos Daskalakis, and S Matthew Weinberg. “An algorithmic characterization of multi-dimensional mechanisms”. In: *Proceedings of the forty-fourth annual ACM symposium on Theory of computing*. 2012, pp. 459–478.
- [3] Yang Cai, Constantinos Daskalakis, and S Matthew Weinberg. “Optimal multi-dimensional mechanism design: Reducing revenue to welfare maximization”. In: *2012 IEEE 53rd Annual Symposium on Foundations of Computer Science*. IEEE. 2012, pp. 130–139.
- [4] Yang Cai, Constantinos Daskalakis, and S Matthew Weinberg. “Understanding incentives: Mechanism design becomes algorithm design”. In: *2013 IEEE 54th Annual Symposium on Foundations of Computer Science*. IEEE. 2013, pp. 618–627.
- [5] L Elisa Celis, Anay Mehrotra, and Nisheeth K Vishnoi. “Toward controlling discrimination in online ad auctions”. In: *International Conference on Machine Learning (ICML)*. 2019.
- [6] Shuchi Chawla and Meena Jagadeesan. “Fairness in ad auctions through inverse proportionality”. In: *arXiv preprint arXiv:2003.13966* (2020).
- [7] Edward H Clarke. “Multipart pricing of public goods”. In: *Public choice* (1971), pp. 17–33.
- [8] Peter Cramton. “Electricity market design”. In: *Oxford Review of Economic Policy* 33.4 (2017), pp. 589–612.
- [9] Peter Cramton. “The FCC spectrum auctions: An early assessment”. In: *Journal of Economics & Management Strategy* 6.3 (1997), pp. 431–495.
- [10] Michael Curry, Ping-yeh Chiang, Tom Goldstein, and John P. Dickerson. “Certifying Strategyproof Auction Networks”. In: *Neural Information Processing Systems (NeurIPS)*. 2020.
- [11] Constantinos Daskalakis, Alan Deckelbaum, and Christos Tzamos. “Mechanism design via optimal transport”. In: *Proceedings of the fourteenth ACM conference on Electronic commerce*. 2013, pp. 269–286.
- [12] Constantinos Daskalakis, Alan Deckelbaum, and Christos Tzamos. “Strong Duality for a Multiple-Good Monopolist”. In: *Econometrica* 85.3 (2017), pp. 735–767.
- [13] Paul Duetting, Zhe Feng, Harikrishna Narasimhan, David C. Parkes, and Sai Srivatsa Ravindranath. “Optimal Auctions through Deep Learning”. In: *International Conference on Machine Learning (ICML)*. 2019.
- [14] Benjamin Edelman, Michael Ostrovsky, and Michael Schwarz. “Internet advertising and the generalized second-price auction: Selling billions of dollars worth of keywords”. In: *American Economic Review* 97.1 (2007), pp. 242–259.
- [15] Zhe Feng, Harikrishna Narasimhan, and David C. Parkes. “Deep Learning for Revenue-Optimal Auctions with Budgets”. In: *International Conference on Autonomous Agents and MultiAgent Systems (AAMAS)*. 2018.
- [16] Sainyam Galhotra, Sandhya Saisubramanian, and Shlomo Zilberstein. “Learning to Generate Fair Clusters from Demonstrations”. In: (Feb. 2021). arXiv: 2102.03977 [stat.ML].
- [17] Robert Geirhos, Jörn-Henrik Jacobsen, Claudio Michaelis, Richard Zemel, Wieland Brendel, Matthias Bethge, and Felix A Wichmann. “Shortcut learning in deep neural networks”. In: *Nature Machine Intelligence* 2.11 (2020), pp. 665–673.
- [18] Yiannis Giannakopoulos and Elias Koutsoupias. “Duality and optimality of auctions for uniform distributions”. In: *Economics and Computation (EC)*. 2014.
- [19] Noah Golowich, Harikrishna Narasimhan, and David C. Parkes. “Deep Learning for Multi-Facility Location Mechanism Design”. In: *International Joint Conference on Artificial Intelligence (IJCAI)*. 2018.

- [20] Theodore Groves. “Incentives in teams”. In: *Econometrica: Journal of the Econometric Society* (1973), pp. 617–631.
- [21] Nima Haghpanah and Jason D. Hartline. “Reverse Mechanism Design”. In: *CoRR* abs/1404.1341 (2014). arXiv: 1404.1341.
- [22] Kenneth Holstein, Jennifer Wortman Vaughan, Hal Daumé III, Miro Dudik, and Hanna Wallach. “Improving fairness in machine learning systems: What do industry practitioners need?” In: *Conference on Human Factors in Computing Systems (CHI)*. 2019.
- [23] Christina Ilvento, Meena Jagadeesan, and Shuchi Chawla. “Multi-Category Fairness in Sponsored Search Auctions”. In: *Conference on Fairness, Accountability, and Transparency (FAccT)*. 2020.
- [24] Ian A Kash and Rafael M Frongillo. “Optimal auctions with restricted allocations”. In: *Economics and Computation (EC)*. 2016.
- [25] Kevin Kuo, Anthony Ostuni, Elizabeth Horishny, Michael J. Curry, Samuel Dooley, Ping-yeh Chiang, Tom Goldstein, and John P. Dickerson. “ProportionNet: Balancing Fairness and Revenue for Auction Design with Deep Learning”. In: *arXiv preprint arXiv:2010.06398* (2020).
- [26] Kevin Leyton-Brown, Paul Milgrom, and Ilya Segal. “Economics and computer science of a radio spectrum reallocation”. In: *Proceedings of the National Academy of Sciences (PNAS)* 114.28 (2017), pp. 7202–7209.
- [27] Jingling Li, Mozhi Zhang, Keyulu Xu, John P. Dickerson, and Jimmy Ba. “Noisy Labels Can Induce Good Representations”. In: *CoRR* abs/2012.12896 (2020). URL: <https://arxiv.org/abs/2012.12896>.
- [28] Alejandro M. Manelli and Daniel R. Vincent. “Bundling as an optimal selling mechanism for a multiple-good monopolist”. In: *J. Econ. Theory* 127.1 (2006), pp. 1–35.
- [29] Roger B Myerson. “Optimal auction design”. In: *Mathematics of Operations Research* 6.1 (1981), pp. 58–73.
- [30] Congressional Budget Office. *The Budget and Economic Outlook: Fiscal Years 2001–2010*. Appendix B. 2000.
- [31] Gregory Pavlov. “Optimal mechanism for selling two goods”. In: *The BE Journal of Theoretical Economics* 11.1 (2011).
- [32] Jad Rahme, Samy Jelassi, Joan Bruna, and S Matthew Weinberg. “A Permutation-Equivariant Neural Network Architecture For Auction Design”. In: (2020).
- [33] Jad Rahme, Samy Jelassi, and S Matthew Weinberg. “Auction learning as a two-player game”. In: *International Conference on Learning Representations (ICLR)*. 2021.
- [34] Jean-Charles Rochet. “A necessary and sufficient condition for rationalizability in a quasi-linear context”. In: *Journal of mathematical Economics* 16.2 (1987), pp. 191–200.
- [35] Debjani Saha, Candice Schumann, Duncan C McElfresh, John P Dickerson, Michelle L Mazurek, and Michael Carl Tschantz. “Measuring Non-Expert Comprehension of Machine Learning Fairness Metrics”. In: *International Conference on Machine Learning (ICML)*. 2020.
- [36] Nripsuta Ani Saxena, Karen Huang, Evan DeFilippis, Goran Radanovic, David C Parkes, and Yang Liu. “How do fairness definitions fare? Examining public attitudes towards algorithmic definitions of fairness”. In: *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society*. 2019, pp. 99–106.
- [37] Weiran Shen, Pingzhong Tang, and Song Zuo. “Automated mechanism design via neural networks”. In: *Proceedings of the 18th International Conference on Autonomous Agents and Multiagent Systems*. 2019, pp. 215–223.
- [38] Megha Srivastava, Hoda Heidari, and Andreas Krause. “Mathematical Notions vs. Human Perception of Fairness: A Descriptive Approach to Fairness for Machine Learning”. In: (Feb. 2019). arXiv: 1902.04783 [cs.LG].
- [39] Andrea Tacchetti, DJ Strouse, Marta Garnelo, Thore Graepel, and Yoram Bachrach. “A neural architecture for designing truthful and efficient auctions”. In: *arXiv preprint arXiv:1907.05181* (2019).
- [40] William Vickrey. “Counterspeculation, auctions, and competitive sealed tenders”. In: *The Journal of finance* 16.1 (1961), pp. 8–37.
- [41] Andrew Chi-Chih Yao. “Dominant-strategy versus bayesian multi-item auctions: Maximum revenue determination and comparison”. In: *Economics and Computation (EC)*. 2017.