

Machine Learning-Powered Combinatorial Clock Auction

Ermis Nikiforos Soumalias^{1,3†}, Jakob Weissteiner^{1,3†}, Jakob Heiss^{2,3}, Sven Seuken^{1,3}

¹University of Zurich

²ETH Zurich

³ETH AI Center

ermis@ifi.uzh.ch, weissteiner@ifi.uzh.ch, jakob.heiss@math.ethz.ch, seuken@ifi.uzh.ch

Abstract

We study the design of *iterative combinatorial auctions* (ICAs). The main challenge in this domain is that the bundle space grows exponentially in the number of items. To address this, several papers have recently proposed machine learning (ML)-based preference elicitation algorithms that aim to elicit only the most important information from bidders. However, from a practical point of view, the main shortcoming of this prior work is that those designs elicit bidders' preferences via *value queries* (i.e., "What is your value for the bundle $\{A, B\}$?"). In most real-world ICA domains, value queries are considered impractical, since they impose an unrealistically high cognitive burden on bidders, which is why they are not used in practice. In this paper, we address this shortcoming by designing an *ML-powered combinatorial clock auction* that elicits information from the bidders only via *demand queries* (i.e., "At prices p , what is your most preferred bundle of items?"). We make two key technical contributions: First, we present a novel method for training an ML model on demand queries. Second, based on those trained ML models, we introduce an efficient method for determining the demand query with the highest clearing potential, for which we also provide a theoretical foundation. We experimentally evaluate our ML-based demand query mechanism in several spectrum auction domains and compare it against the most established real-world ICA: the *combinatorial clock auction* (CCA). Our mechanism significantly outperforms the CCA in terms of efficiency in all domains, it achieves higher efficiency in a significantly reduced number of rounds, and, using linear prices, it exhibits vastly higher clearing potential. Thus, with this paper we bridge the gap between research and practice and propose the first practical ML-powered ICA.

1 Introduction

Combinatorial auctions (CAs) are used to allocate multiple items among several bidders who may view those items as complements or substitutes. In a CA, bidders are allowed to submit bids over *bundles* of items. CAs have enjoyed widespread adoption in practice, with their applications ranging from allocating spectrum licences (Cramton 2013) to TV ad slots (Goetzendorff et al. 2015) and airport landing/take-off slots (Rassenti, Smith, and Bulfin 1982).

*This paper is the full version of Soumalias et al. (2024) published at AAAI'24 including the appendix.

†These authors contributed equally.

One of the key challenges in CAs is that the bundle space grows exponentially in the number of items, making it infeasible for bidders to report their full value function in all but the smallest domains. Moreover, Nisan and Segal (2006) showed that for general value functions, CAs require an exponential number of bids in order to achieve full efficiency in the worst case. Thus, practical CA designs cannot provide efficiency guarantees in real world settings with more than a modest number of items. Instead, the focus has shifted towards *iterative combinatorial auctions* (ICAs), where bidders interact with the auctioneer over a series of rounds, providing a limited amount of information, and the aim of the auctioneer is to find a highly efficient allocation.

The most established mechanism following this interaction paradigm is the *combinatorial clock auction* (CCA) (Ausubel, Cramton, and Milgrom 2006). The CCA has been used extensively for spectrum allocation, generating over \$20 Billion in revenue between 2012 and 2014 alone (Ausubel and Baranov 2017). Speed of convergence is a critical consideration for any ICA since each round can entail costly computations and business modelling for the bidders (Kwasnica et al. 2005; Milgrom and Segal 2017; Bichler, Hao, and Adomavicius 2017). Large spectrum auctions following the CCA format can take more than 100 bidding rounds. In order to decrease the number of rounds, many CAs in practice use aggressive price update rules (e.g., increasing prices by up to 10% each round), which can harm efficiency (Ausubel and Baranov 2017). Thus, it remains a challenging problem to design a practical ICA that elicits information via demand queries, is efficient, and converges in a small number of rounds. Specifically, given the value of resources allocated in such real-world ICAs, increasing their efficiency by even one percentage point already translates into monetary gains of hundreds of millions of dollars.

1.1 ML-Powered Preference Elicitation

To address this challenge, researchers have proposed various ways of using machine learning (ML) to improve the efficiency of CAs. The seminal works by Blum et al. (2004) and Lahaie and Parkes (2004) were the first to frame preference elicitation in CAs as a learning problem. In the same strand of research, Brero, Lubin, and Seuken (2018, 2021), Weissteiner and Seuken (2020) and Weissteiner et al. (2022b) proposed ML-powered ICAs. At the heart of those

approaches lies an ML-powered preference elicitation algorithm that uses an ML model to approximate each bidder’s value function and to generate the next value query, which in turn refines that bidder’s model. Weissteiner et al. (2022a) designed a special network architecture for this framework while Weissteiner et al. (2023) incorporated a notion of uncertainty (Heiss et al. 2022) into the framework, further increasing its efficiency. Despite their great efficiency gains compared to traditional CA designs, those approaches suffer from one common limitation: they fundamentally rely on value queries of the form “What is your value for bundle $\{A, B\}$ ”. Prior research in auction design has identified demand queries (DQs) as the best way to run an auction (Cramton 2013). Their advantages compared to value queries include elimination of tacit collusion and bid signaling, as well as simplified bidder decision-making that keeps the bidders focused on what is most relevant: the relationship between prices and aggregate demand. Additionally, value queries are cognitively more complex, and thus typically impractical for real-world ICAs. For these reasons, DQs are the most prominent interaction paradigm for auctions in practice.

Despite the prominence of DQs in real-world applications, the only prior work on ML-based DQs that we are aware of is that of Brero and Lahaie (2018) and Brero, Lahaie, and Seuken (2019), who proposed integrating ML in a price-based ICA to generate the next price vector in order to achieve faster convergence. Similar to our design, in these works the auctioneer maintains a model of each agent’s value function, which are updated as the agents bid in the auction and reveal more information about their values. Then, those models are used in each round to compute new prices and drive the bidding process. Unlike our approach, the design of this prior work focuses solely on clearing potential, as the authors do not report efficiency results. Additionally, their design suffers from some significant limitations: (i) it does not exploit any notion of similarity between bundles that contain overlapping items, (ii) it only incorporates a fraction of the information revealed by the agents’ bidding. Specifically, it only makes use of the fact that for the bundle an agent bids on, her value for that bundle must be larger than its price, and (iii) their approach is computationally intractable already in medium-sized auction domains, as their price update rule requires a large number l of posterior samples for *expectation maximization* and then solving a linear program whose number of constraints for each bidder is proportional to l times the number of bids by that agent. These limitations are significant, as they can lead to large efficiency decreases in complex combinatorial domains. Moreover, their design cannot be easily modified to alleviate these limitations. In contrast, our approach effectively addresses all of these limitations.

1.2 Our Contributions

In this paper, we address the main shortcomings of prior work by designing an ML-powered combinatorial clock auction. Our auction elicits information from bidders via *demand queries (DQs)* instead of value queries, while simultaneously, unlike prior work on ML-based DQs, being computationally feasible for large domains and incorporating

the *complete information* the demand query observations provide into the training of our ML models. Concretely, we use *Monotone-Value Neural Networks (MVNNs)* (Weissteiner et al. 2022a) as ML models, which are tailored to model monotone and non-linear combinatorial value functions in CAs.

The main two technical challenges are (i) training those MVNNs only on *demand query* observations and (ii) efficiently determining the next demand query that is most likely to clear the market based on the trained MVNNs. In detail, we make the following contributions:

1. We first propose an adjusted MVNN architecture, which we call *multiset MVNNs (mMVNNs)*. mMVNNs can be used more generally in *multiset domains* (i.e., if multiple indistinguishable copies of the same good exist) (Section 3.1) and we prove the universality property of mMVNNs in such multiset domains (Theorem 1).
2. We introduce a novel method for training *any* MIP-formalizable and gradient descent (GD)-compatible ML model (e.g., mMVNNs) on *demand query* observations (Section 3.2).¹ Unlike prior work, our training method provably makes use of the *complete* information provided by the demand query observations.
3. We introduce an efficient method for determining the price vector that is most likely to clear the market based on the trained ML models (Section 4). For this, we derive a simple and intuitive price update rule that results from performing GD on an objective function which is minimized exactly at clearing prices (Theorem 3).
4. Based on Items 2 and 3, we propose a practical ML-powered clock auction (Section 5).
5. We experimentally show that compared to the CCA, our ML-powered clock auction can achieve substantially higher efficiency on the order of 9% points. Furthermore, using linear prices, our ML-powered clock auction exhibits significantly higher clearing potential compared to the CCA (Section 6).

GitHub Our source code is publicly available on GitHub at <https://github.com/marketdesignresearch/ML-CCA>.

1.3 Further Related Work

In the field of *automated mechanism design*, Dütting et al. (2015, 2019), Golowich, Narasimhan, and Parkes (2018) and Narasimhan, Agarwal, and Parkes (2016) used ML to learn new mechanisms from data, while Cole and Roughgarden (2014); Morgenstern and Roughgarden (2015) and Balcan, Sandholm, and Vitercik (2023) bounded the sample complexity of learning approximately optimal mechanisms. In contrast to this prior work, our design incorporates an ML algorithm into the mechanism itself, i.e., the ML algorithm is part of the mechanism. Lahaie and Lubin (2019) suggest an adaptive price update rule that increases price expressivity as the rounds progress in order to improve efficiency and speed of convergence. Unlike that work, we aim to improve preference elicitation while still using linear prices. Preference elicitation is a key market design challenge outside of

¹Namely, this includes neural networks with any piecewise linear activation function.

CAs too. [Soumalias et al. \(2023\)](#) introduce an ML-powered mechanism for course allocation that improves preference elicitation by asking students comparison queries.

1.4 Practical Considerations and Incentives

Our ML-powered clock phase can be viewed as an alternative to the clock phase of the CCA. In a real-world application, many other considerations (beyond the price update rule) are also important. For example, the careful design of *activity rules* is vital to induce truthful bidding in the clock phase of the CCA ([Ausubel and Baranov 2017](#)). The payment rule used in the supplementary round is also important, and it has been argued that the use of the VCG-nearest payment rule, while not strategy-proof, induces good incentives in practice ([Cramton 2013](#)). Similar to the clock phase of the CCA, our ML-powered clock phase is not strategyproof. If our design were to be fielded in a real-world environment, we envision that one would combine it with carefully designed activity and payment rules in order to induce good incentives. Thus, we consider the incentive problem orthogonal to the price update problem and in the rest of the paper, we follow prior work ([Brero, Lahaie, and Seuken 2019](#); [Parkes and Ungar 2000](#)) and assume that bidders follow myopic best-response (truthful) bidding throughout all auction mechanisms tested.

2 Preliminaries

2.1 Formal Model for ICAs

We consider *multiset* CA domains with a set $N = \{1, \dots, n\}$ of bidders and a set $M = \{1, \dots, m\}$ of distinct items with corresponding *capacities*, i.e., number of available copies, $c = (c_1, \dots, c_m) \in \mathbb{N}^m$. We denote by $x \in \mathcal{X} = \{0, \dots, c_1\} \times \dots \times \{0, \dots, c_m\}$ a bundle of items represented as a positive integer vector, where $x_j = k$ iff item $j \in M$ is contained k -times in x . The bidders' true preferences over bundles are represented by their (private) value functions $v_i : \mathcal{X} \rightarrow \mathbb{R}_{\geq 0}$, $i \in N$, i.e., $v_i(x)$ represents bidder i 's true value for bundle $x \in \mathcal{X}$. We collect the value functions v_i in the vector $v = (v_i)_{i \in N}$. By $a = (a_1, \dots, a_n) \in \mathcal{X}^n$ we denote an allocation of bundles to bidders, where a_i is the bundle bidder i obtains. We denote the set of *feasible* allocations by $\mathcal{F} = \{a \in \mathcal{X}^n : \sum_{i \in N} a_{ij} \leq c_j, \forall j \in M\}$. We assume that bidders have quasilinear utility functions u_i of the form $u_i(a_i) = v_i(a_i) - \pi_i$ where v_i can be highly non-linear and $\pi_i \in \mathbb{R}_{\geq 0}$ denotes the bidder's payment. This implies that the (true) *social welfare* $V(a)$ of an allocation a is equal to the sum of all bidders' values $\sum_{i \in N} v_i(a_i)$. We let $a^* \in \operatorname{argmax}_{a \in \mathcal{F}} V(a)$ denote a social-welfare maximizing, i.e., *efficient*, allocation. The *efficiency* of any allocation $a \in \mathcal{F}$ is determined as $V(a)/V(a^*)$.

An ICA *mechanism* defines how the bidders interact with the auctioneer and how the allocation and payments are determined. In this paper, we consider ICAs that iteratively ask bidders *linear demand queries*. In such a query, the auctioneer presents a vector of item prices $p \in \mathbb{R}_{\geq 0}^m$ and each bidder i responds with her utility-maximizing bundle, i.e.,

$$x_i^*(p) \in \operatorname{argmax}_{x \in \mathcal{X}} \{v_i(x) - \langle p, x \rangle\} \quad i \in N, \quad (1)$$

where $\langle \cdot, \cdot \rangle$ denotes the Euclidean scalar product in $\mathbb{R}^m$.

Even though our approach could conceptually incorporate any kind of (non-linear) price function $p : \mathcal{X} \rightarrow \mathbb{R}_{\geq 0}$, our concrete implementation will only use linear prices (i.e., prices over items). Linear prices are most established in practice since they are intuitive and simple for the bidders to understand (e.g., ([Ausubel, Cramton, and Milgrom 2006](#))).

For bidder $i \in N$, we denote a set of $K \in \mathbb{N}$ such elicited utility-maximizing bundles and price pairs as $R_i = \{(x_i^*(p^r), p^r)\}_{r=1}^K$. Let $R = (R_1, \dots, R_n)$ be the tuple of elicited demand query data from all bidders. The ICA's (inferred) optimal feasible allocation $a^*(R) \in \mathcal{F}$ and payments $\pi_i := \pi_i(R) \in \mathbb{R}_+^n$ are computed based on the elicited reports R only. Concretely, $a_R^* \in \mathcal{F}$ is defined as

$$a^*(R) \in \operatorname{argmax}_{\substack{\mathcal{F} \ni (x_i^*(p^{r_i}))_{i=1}^n \\ : r_i \in \{1, \dots, K\} \forall i \in N}} \sum_{i \in N} \langle p^{r_i}, x_i^*(p^{r_i}) \rangle. \quad (2)$$

In words, a bidder's response to a demand query provides a lower bound on that bidder's value for the bundle she requested. That lower bound is equal to the bundle's price in the round the bundle was requested. The ICA's optimal (inferred) feasible allocation $a^*(R) \in \mathcal{F}$ is the one that maximizes the corresponding lower bound on social welfare, based on all elicited demand query data R from the bidders. As payment rule $\pi_i(R)$ one could use any reasonable choice (e.g., VCG payments, see Appendix A). As the auctioneer can only ask a limited number of demand queries $|R_i| \leq Q^{\max}$ (e.g., $Q^{\max} = 100$), an ICA needs a practically feasible and smart preference elicitation algorithm.

2.2 The Combinatorial Clock Auction (CCA)

We consider the CCA ([Ausubel, Cramton, and Milgrom 2006](#)) as the main benchmark auction. The CCA consists of two phases. The initial *clock phase* proceeds in rounds. In each round, the auctioneer presents anonymous item prices $p \in \mathbb{R}_{\geq 0}^m$, and each bidder is asked to respond to a *demand query*, declaring her utility-maximizing bundle at p . The clock phase of the CCA is parametrized by the *reserve prices* employed in its first round, and the way prices are updated. An item j is *over-demanded* at prices p , if, for those prices, its total demand based on the bidders' responses to the demand query exceeds its capacity, i.e., $\sum_{i \in N} (x_i^*(p))_j > c_j$. The most common price update rule is to increase the price of all over-demanded items by a fixed percentage, which we set to 5% for our experiments, as in many real-world applications (e.g., ([Industry Canada 2013](#))).

The second phase of the CCA is the *supplementary round*. In this phase, each bidder can submit a finite number of additional bids for bundles of items, which are called *push bids*. Then, the final allocation is determined based on the combined set of all inferred bids of the clock phase, plus all submitted push bids of the supplementary round. This design aims to combine good price discovery in the clock phase with good expressiveness in the supplementary round. In simulations, the supplementary round is parametrized by the assumed bidder behaviour in this phase, i.e., which bundle-value pairs they choose to report. As in ([Brero, Lubin, and](#)

Algorithm 1: TRAINONDQS

Input : Demand query data $R_i = \{(x_i^*(p^r), p^r)\}_{r=1}^K$,
Epochs $T \in \mathbb{N}$, Learning Rate $\gamma > 0$.

- 1 $\theta_0 \leftarrow \text{init mMVNN} \triangleright \text{Weissteiner et al. (2023, S.3.2)}$
- 2 **for** $t = 0$ **to** $T - 1$ **do**
- 3 **for** $r = 1$ **to** K **do** $\triangleright \text{Demand responses for prices}$
- 4 Solve $\hat{x}_i^*(p^r) \in \arg\max_{x \in \mathcal{X}} \mathcal{M}_i^{\theta_t}(x) - \langle p^r, x \rangle$
- 5 **if** $\hat{x}_i^*(p^r) \neq x_i^*(p^r)$ **then** $\triangleright \text{mMVNN is wrong}$
- 6 $L(\theta_t) \leftarrow (\mathcal{M}_i^{\theta_t}(\hat{x}_i^*(p^r)) - \langle p^r, \hat{x}_i^*(p^r) \rangle) -$
 $(\mathcal{M}_i^{\theta_t}(x_i^*(p^r)) - \langle p^r, x_i^*(p^r) \rangle) \triangleright \text{Add}$
 $\text{predicted utility difference to loss}$
- 7 $\theta_{t+1} \leftarrow \theta_t - \gamma(\nabla_{\theta} L(\theta))_{\theta=\theta_t} \triangleright \text{SGD step}$
- 8 **return** Trained parameters θ_T of the mMVNN $\mathcal{M}_i^{\theta_T}$

Seuken 2021), we consider the following heuristics when simulating bidder behaviour:

- **Clock Bids:** Corresponds to having no supplementary round. Thus, the final allocation is determined based only on the inferred bids of the clock phase (Equation (2)).
- **Raised Clock Bids:** The bidders also provide their true value for all bundles they bid on during the clock phase.
- **Profit Max:** Bidders provide their true value for all bundles that they bid on in the clock phase, and additionally submit their true value for the $Q^{\text{P-Max}}$ bundles earning them the highest utility at the prices of the final clock phase.

3 Training on Demand Query Observations

In this section, we first propose a new version of MVNNs that are applicable to multiset domains $\mathcal{X}$ and extend the universality proof of classical MVNNs. Finally, we present our demand-query training algorithm.

3.1 Multiset MVNNs

MVNNs (Weissteiner et al. 2022a) are a recently introduced class of NNs specifically designed to represent *monotone combinatorial* valuations. We introduce an adapted version of MVNNs, which we call *multiset MVNNs* (mMVNNs). Compared to MVNNs, mMVNNs have an added linear normalization layer D after the input layer. We add this normalization since the input (i.e., a bundle) $x \in \mathcal{X}$ is a positive integer vector instead of a binary vector as in the classic case of indivisible items with capacities $c_j = 1$ for all $j \in M$. This normalization ensures that $Dx \in [0, 1]$ and thus we can use the weight initialization scheme from (Weissteiner et al. 2023). Unlike MVNNs, mMVNNs incorporate at a structural level the prior information that some items are identical and consequently significantly reduce the dimensionality of the input space. This improves the sample efficiency of mMVNNs, which is especially important in applications with a limited number of samples such as auctions. For more details on mMVNNs and their advantages, please see Appendix B.

Definition 1 (Multiset MVNN). An mMVNN $\mathcal{M}_i^{\theta} : \mathcal{X} \rightarrow \mathbb{R}_{\geq 0}$ for bidder $i \in N$ is defined as

$$\mathcal{M}_i^{\theta}(x) := W^{i, K_i} \varphi_{0, t^{i, K_i}-1} \left(\dots \varphi_{0, t^{i, 1}}(W^{i, 1}(Dx) + b^{i, 1}) \dots \right) \quad (3)$$

- $K_i + 2 \in \mathbb{N}$ is the number of layers (K_i hidden layers),
- $\{\varphi_{0, t^{i, k}}\}_{k=1}^{K_i-1}$ are the MVNN-specific activation functions with cutoff $t^{i, k} > 0$, called bounded ReLU (bReLU):

$$\varphi_{0, t^{i, k}}(\cdot) := \min(t^{i, k}, \max(0, \cdot)) \quad (4)$$

- $W^i := (W^{i, k})_{k=1}^{K_i}$ with $W^{i, k} \geq 0$ and $b^i := (b^{i, k})_{k=1}^{K_i-1}$ with $b^{i, k} \leq 0$ are the non-negative weights and non-positive biases of dimensions $d^{i, k} \times d^{i, k-1}$ and $d^{i, k}$, whose parameters are stored in $\theta = (W^i, b^i)$.
- $D := \text{diag}(1/c_1, \dots, 1/c_m)$ is the linear normalization layer that ensures $Dx \in [0, 1]$ and is not trainable.

In Theorem 1, we extend the proof from Weissteiner et al. (2022a) and show that mMVNNs are also universal in the set of monotone value functions defined on a multiset domain $\mathcal{X}$. For this, we first define the following properties:

- (M) **Monotonicity** (“more items weakly increase value”):
For $a, b \in \mathcal{X}$: if $a \leq b$, i.e. $\forall k \in M : a_k \leq b_k$, it holds that $v_i(a) \leq v_i(b)$,
- (N) **Normalization** (“no value for empty bundle”):
 $v_i(\emptyset) = v_i((0, \dots, 0)) := 0$,

These properties are common assumptions and are satisfied in many market domains. We can now present the following universality result:

Theorem 1 (Multiset Universality). Any value function $v_i : \mathcal{X} \rightarrow \mathbb{R}_+$ that satisfies (M) and (N) can be represented exactly as an mMVNN $\mathcal{M}_i^{\theta}$ from Definition 1, i.e., for $\mathcal{V} := \{v_i : \mathcal{X} \rightarrow \mathbb{R}_+ \mid \text{satisfy (M) and (N)}\}$ it holds that

$$\mathcal{V} = \left\{ \mathcal{M}_i^{(W^i, b^i)} : W^i \geq 0, b^i \leq 0 \right\}. \quad (5)$$

Proof. Please, see Appendix B.2 for the proof. $\square$

Furthermore, we can formulate maximization over mMVNNs, i.e., $\max_{x \in \mathcal{X}} \mathcal{M}_i^{\theta}(x) - \langle p, x \rangle$, as a *mixed integer linear program* (MILP) analogously to Weissteiner et al. (2022a), which will be key for our ML-powered clock phase.

3.2 Training Algorithm

In Algorithm 1, we describe how we train, for each bidder $i \in N$, a distinct mMVNN $\mathcal{M}_i^{\theta}$ on demand query data R_i . Our design choices regarding this training algorithm are motivated by the information that responses to demand queries provide. According to myopic best response bidding, at each round r , bidder i reports a utility-maximizing bundle $x_i^*(p^r) \in \mathcal{X}$ at current prices p^r . Formally, for all $x \in \mathcal{X}$:

$$v_i(x_i^*(p^r)) - \langle p^r, x_i^*(p^r) \rangle \geq v_i(x) - \langle p^r, x \rangle. \quad (6)$$

Notice that for any epoch t and round r , the loss $L(\theta_t)$ for that round calculated in Lines 4 to 6 is always non-negative, and can only be zero if the mMVNN $\mathcal{M}_i^{\theta_t}$ (instead of v_i) satisfies Equation (6). Thus, the loss for an epoch is zero iff the mMVNN $\mathcal{M}_i^{\theta_t}$ satisfies Equation (6) for all rounds, and in that case the model has captured the full information provided by the demand query responses R_i of that bidder. Finally, note that Algorithm 1 can be applied to any MILP-formalizable ML model whose parameters can be efficiently updated via GD, such as MVNNs or ReLU-NNs.

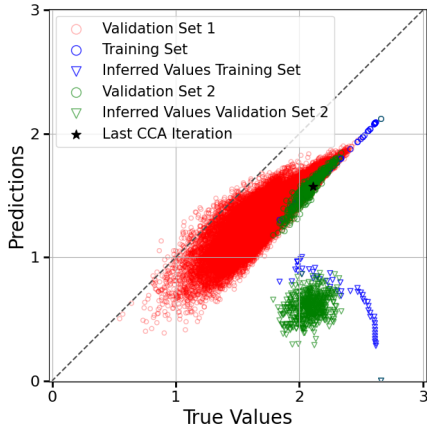


Figure 1: Scaled prediction vs. true plot of a trained mMVNN via Algorithm 1 for the national bidder in the MRVM domain (see Section 6).

In Figure 1, we present a prediction vs. true plot of an mMVNN, which we trained via Algorithm 1. We present the training set of 50 demand query data points R_i in blue circles, where the prices $\{p^r\}_{r=1}^{50}$ are generated according to the same rule as in CCA. Additionally, we mark the bundle $x^{\text{CCA}} \in \mathcal{X}$ from this last CCA iteration (i.e., the one resulting from p^{50}) with a black star. Moreover, we present two different validation sets on which we evaluate mMVNN configurations in our hyperparameter optimization (HPO): *Validation set 1* (red circles), which are 50,000 uniformly at random sampled bundles $x \in \mathcal{X}$, and *validation set 2* (green circles), where we first sample 500 price vectors $\{p^r\}_{r=1}^{500}$ where the price of each item is drawn uniformly at random from the range of 0 to 3 times the average maximum value of an agent of that type for a single item, and then determine utility-maximizing bundles $x_i^*(p^r)$ (w.r.t. v_i) at those prices (cp. Equation (1)). While validation set 1 measures generalization performance in a classic sense over the whole bundle space, validation set 2 focuses on utility-maximizing bundles. We additionally demonstrate the inferred values of the bundles of the training set and validation set 2 using triangles of the same colour, i.e., $\{\langle p^r, x_i^*(p^r) \rangle\}_{r=1}^{50/500}$. These triangles highlight the only cardinal information that our mMVNNs have access to during training and are a lower bound of the true value. In Figure 1, we see that our mMVNN is able to learn at the training points (blue circles) the true value functions almost perfectly up to a constant shift κ , i.e., $\mathcal{M}_i^{\theta_T}(x) \approx v_i(x) + \kappa$. This is true even though the corresponding inferred values (blue triangles) are very far off from the true values $\langle p^r, x_i^*(p^r) \rangle \ll v_i(x_i^*(p^r))$. Moreover, the mMVNN generalizes well (up to the constant shift κ) on validation sets 1 and 2. Overall, this shows that Algorithm 1 indeed leads to mMVNNs $\mathcal{M}_i^{\theta_T}$ which are a good approximation of $v_i + \kappa$. Note that learning the true value function up to a constant shift suffices for our proposed demand query generation procedure presented in Section 4.

4 ML-powered Demand Query Generation

In this section, we show how we generate ML-powered demand queries and provide the theoretical foundation for our approach by extending a well-known connection between *clearing prices*, *efficiency* and a *clearing objective function*. First, we define indirect utility, revenue and clearing prices.

Definition 2 (Indirect Utility and Revenue). *For linear prices $p \in \mathbb{R}_{\geq 0}^m$, a bidder's indirect utility U and the seller's indirect revenue R are defined as*

$$U(p, v_i) := \max_{x \in \mathcal{X}} \{v_i(x) - \langle p, x \rangle\} \text{ and} \quad (7)$$

$$R(p) := \max_{a \in \mathcal{F}} \left\{ \sum_{i \in N} \langle p, a_i \rangle \right\} = \frac{1}{\sum_{j \in M} c_j} \sum_{j \in M} c_j p_j, \quad (8)$$

i.e., at prices p , Equations (7) and (8) are the maximum utility a bidder can achieve for all $x \in \mathcal{X}$ and the maximum revenue the seller can achieve among all feasible allocations.

Definition 3 (Clearing Prices). *Prices $p \in \mathbb{R}_{\geq 0}^m$ are clearing prices if there exists an allocation $a(p) \in \mathcal{F}$ such that*

1. *for each bidder i , the bundle $a_i(p)$ maximizes her utility, i.e., $v_i(a_i(p)) - \langle p, a_i(p) \rangle = U(p, v_i), \forall i \in N$, and*
2. *the allocation $a(p) \in \mathcal{F}$ maximizes the seller's revenue, i.e., $\sum_{i \in N} \langle p, a_i(p) \rangle = R(p)$.²*

Next, we provide an important connection between *clearing prices*, *efficiency* and a *clearing objective W* . Theorem 2 extends Bikhchandani and Ostroy (2002, Theorem 3.1).

Theorem 2. *Consider the notation from Definitions 2 and 3 and the objective function $W(p, v) := R(p) + \sum_{i \in N} U(p, v_i)$. Then it holds that, if a linear clearing price vector exists, every price vector*

$$p' \in \operatorname{argmin}_{\tilde{p} \in \mathbb{R}_{\geq 0}^m} W(\tilde{p}, v) \quad (9a)$$

$$\text{such that } (x_i^*(\tilde{p}))_{i \in N} \in \mathcal{F} \quad (9b)$$

is a clearing price vector and the corresponding allocation $a(p') \in \mathcal{F}$ is efficient.³

Proof. Please, see Appendix C.1 for the proof. $\square$

Theorem 2 does not claim the existence of *linear clearing prices* (LCPs) $p \in \mathbb{R}_{\geq 0}^m$. For general value functions v , LCPs may not exist (Bikhchandani and Ostroy 2002). However, in the case that LCPs do exist, Theorem 2 shows that *all* minimizers of (9) are LCPs and their corresponding allocation is efficient. This is at the core of our ML-powered demand query generation approach, which we discuss next.

The key idea to generate ML-powered demand queries is as follows: As an approximation for the true value function

²For linear prices, this maximum is achieved by selling every item, i.e., $\forall j \in M : \sum_{i \in N} (a_i)_j = c_j$ (see Appendix C.2).

³More precisely, constraint (9b) should be reformulated as

$$\exists (x_i^*(\tilde{p}))_{i \in N} \in \bigtimes_{i \in N} \mathcal{X}_i^*(\tilde{p}) : (x_i^*(\tilde{p}))_{i \in N} \in \mathcal{F},$$

where $\mathcal{X}_i^*(\tilde{p}) := \operatorname{argmax}_{x \in \mathcal{X}} \{v_i(x) - \langle \tilde{p}, x \rangle\}$, since in theory, $x_i^*(\tilde{p})$ does not always have to be unique.

v_i , we use for each bidder a distinct mMVNN $\mathcal{M}_i^\theta : \mathcal{X} \rightarrow \mathbb{R}_{\geq 0}$ that has been trained on the bidder's elicited demand query data R_i (see Section 3). Motivated by Theorem 2, we then try to find the demand query $p \in \mathbb{R}_{\geq 0}^m$ minimizing $W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ subject to the feasibility constraint (9b). This way, we find demand queries $p \in \mathbb{R}_{\geq 0}^m$ which, given the already observed demand responses R , have high clearing potential. Note that unlike the CCA, this process does not result in monotone prices.⁴

Remark 1 (Constraint (9b)). *An important economic insight is that minimizing $W(\cdot, (\mathcal{M}_i^\theta)_{i=1}^n)$ is optimal, when LCPs exist (also without constraint (9b) as shown in Lemma 2 in Appendix C.1). If however LCPs do not exist, it is favourable to minimize W under the constraint of having no predicted over-demand for any items (see Appendix D.9 for an empirical comparison of minimizing W with and without constraint (9b)). This is because in case the market does not clear, our ML-CCA (see Section 5), just like the CCA, will have to combine the clock bids of the agents to produce a feasible allocation with the highest inferred social welfare according to Equation (2). See Appendix D.6 for details.*

Note that (9) is a hard, bi-level optimization problem. We minimize (9) via gradient descent (GD), since Theorem 3 gives us the gradient and convexity of $W(\cdot, (\mathcal{M}_i^\theta)_{i=1}^n)$.

Theorem 3. *Let $(\mathcal{M}_i^\theta)_{i=1}^n$ be a tuple of trained mMVNNs and let $\hat{x}_i^*(p) \in \arg\max_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\}$ denote each bidder's predicted utility maximizing bundle w.r.t. $\mathcal{M}_i^\theta$. Then it holds that $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ is convex, Lipschitz-continuous and a.e. differentiable. Moreover,*

$$c - \sum_{i \in N} \hat{x}_i^*(p) \in \nabla_p^{\text{sub}} W(p, (\mathcal{M}_i^\theta)_{i=1}^n) \quad (10)$$

is always a sub-gradient and a.e. a classical gradient.

Proof. In Appendix C.2 we provide the full proof. Concretely, Lemmas 3 and 4 prove the Lipschitz-continuity and the convexity. In the following, we provide a sketch of how the (sub-)gradients are derived. First, since $\mathcal{X}$ is finite, it is intuitive that $\hat{x}_i^*(p)$ is a piece-wise constant function and thus $\partial_p \hat{x}_i^*(p) \stackrel{\text{a.e.}}{=} 0$ (as intuitively argued by Pogančić et al. (2020) and proven by us in Lemma 6). Then we can compute

the gradient a.e. as if $\hat{x}_i^*(p)$ was a constant:

$$\begin{aligned} \nabla_p W(p, (\mathcal{M}_i^\theta)_{i=1}^n) &= \nabla_p \left(R(p) + \sum_{i \in N} U(p, \mathcal{M}_i^\theta) \right) \\ &= \nabla_p \left(\sum_{j \in M} c_j p_j + \sum_{i \in N} (\mathcal{M}_i^\theta(\hat{x}_i^*(p)) - \langle p, \hat{x}_i^*(p) \rangle) \right) \\ &\stackrel{\text{a.e.}}{=} c + \sum_{i \in N} (0 - \nabla_p \langle p, \hat{x}_i^*(p) \rangle) = c - \sum_{i \in N} \hat{x}_i^*(p). \end{aligned}$$

For a mathematically rigorous derivation of sub-gradients and a.e. differentiability see Lemmas 5 and 6. $\square$

With Theorem 3, we obtain the following update rule of classical GD $p_j^{\text{new}} \stackrel{\text{a.e.}}{=} p_j - \gamma(c_j - \sum_{i \in N} (\hat{x}_i^*(p))_j)$, $\forall j \in M$. Interestingly, this equation has an intuitive economic interpretation. If the j^{th} item is over/under-demanded based on the predicted utility-maximizing bundles $\hat{x}_i^*(p)$, then its new price p_j^{new} is increased/decreased by the learning rate times its over/under-demand. However, to enforce constraint (9b) in GD, we asymmetrically increase the prices $1 + \mu \in \mathbb{R}_{\geq 0}$ times more in case of over-demand than we decrease them in case of under-demand. This leads to our final update rule (see Item 1 in Appendix D.6 for more details):

$$p_j^{\text{new}} \stackrel{\text{a.e.}}{=} p_j - \tilde{\gamma}_j(c_j - \sum_{i \in N} (\hat{x}_i^*(p))_j), \quad \forall j \in M, \quad (11a)$$

$$\tilde{\gamma}_j := \begin{cases} \gamma \cdot (1 + \mu) & , c_j < \sum_{i \in N} (\hat{x}_i^*(p))_j \\ \gamma & , \text{else} \end{cases} \quad (11b)$$

To turn this soft constraint into a hard constraint, we increase this asymmetry via μ iteratively until we achieve feasibility and in the end we select the GD step with the lowest W value *within* those steps that were feasible. Based on the final update rule from (11), we propose NEXTPRICE (Algorithm 3 in Appendix D.6), an algorithm that generates demand queries with high clearing potential, which additionally induce utility-maximizing bundles that are predicted to be feasible (see Appendix D.6 for all details).

5 ML-powered Combinatorial Clock Auction

In this section, we describe our *ML-powered combinatorial clock auction (ML-CCA)*, which is based on our proposed new training algorithm from Section 3 as well as our new demand query generation procedure from Section 4.

We present ML-CCA in Algorithm 2. In Lines 2 to 5, we draw the first Q^{init} price vectors using some initial demand query method F_{init} and receive the bidders' demand responses to those price vectors. Concretely, in Section 6, we report results using the same price update rule as the CCA for F_{init} . In each of the next up to $Q^{\text{max}} - Q^{\text{init}}$ ML-powered rounds, we first train, for each bidder, an mMVNN on her demand responses using Algorithm 1 (Line 8). Next, in Line 9, we call NEXTPRICE to generate the next demand query p based on the agents' trained mMVNNs (see Section 4). If, based on the agents' responses to the demand query (Line 11), our algorithm has found market-clearing

⁴We see no reason why non-monotone prices would introduce additional complexity for the bidders. With our approach, the prices quickly converge to the final prices, and then only change very little, as shown in Figures 5 to 8 of Appendix D.7. For this reason, one could even argue that round-over-round optimizations for the bidders may be *easier* in our auction: given that prices are close to each other round-over-round, the optimal bundle from the last round is still close to optimal (in terms of utility) in the next round.

	GSVM				LSVM				SRVM				MRVM			
MECHANISM	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA	98.23	98.93	100.00	56	91.64	96.39	99.95	26	99.59	99.93	100.00	13	93.04	93.31	93.68	0
CCA	90.40	93.59	100.00	3	82.56	91.60	99.76	0	99.63	99.81	100.00	8	92.44	92.62	93.18	0

Table 1: ML-CCA vs CCA. Shown are averages over a test set of 100 synthetic CA instances of the following metrics: efficiency in % for clock bids (E_{CLOCK}), raised clock bids (E_{RAISE}) and raised clock bids plus 100 profit-max bids (E_{PROFIT}) and percentage of instances where clearing prices were found (CLEAR). Winners based on a paired t-test with $\alpha = 5\%$ are marked in grey.

Algorithm 2: ML-CCA($Q^{\text{init}}, Q^{\text{max}}, F_{\text{init}}$)

Parameters: $Q^{\text{init}}, Q^{\text{max}}$ with $Q^{\text{init}} \leq Q^{\text{max}}$ and F_{init}

```

1  $R \leftarrow (\{\})_{i=1}^N$ 
2 for  $r = 1, \dots, Q^{\text{init}}$  do ▷ Draw  $Q^{\text{init}}$  initial prices
3    $p^r \leftarrow F_{\text{init}}(R)$ 
4   foreach  $i \in N$  do ▷ Initial demand query responses
5      $R_i \leftarrow R_i \cup \{(x_i^*(p^r), p^r)\}$ 
6 for  $r = Q^{\text{init}} + 1, \dots, Q^{\text{max}}$  do ▷ ML-powered rounds
7   foreach  $i \in N$  do
8      $\mathcal{M}_i^r \leftarrow \text{TRAINONDQS}(R_i)$  ▷ Algorithm 1
9      $p^r \leftarrow \text{NEXTPRICE}((\mathcal{M}_i^r)_{i=1}^n)$  ▷ Algorithm 3
10    foreach  $i \in N$  do ▷ Demand query responses for  $p^r$ 
11       $R_i \leftarrow R_i \cup \{(x_i^*(p^r), p^r)\}$ 
12    if  $\sum_{i=1}^n (x_i^*(p^r))_j = c_j \forall j \in M$  then ▷ Market-clearing
13      Set final allocation  $a^*(R) \leftarrow (x_i^*(p^r))_{i=1}^n$ 
14      Calculate payments  $\pi(R) \leftarrow (\pi_i(R))_{i=1}^n$ 
15      return  $a^*(R)$  and  $\pi(R)$ 
16 foreach  $i \in N$  do
17    $R_i \leftarrow R_i \cup B_i$  ▷ Optional Push bids
18 Calculate final allocation  $a^*(R)$  as in Equation (2)
19 Calculate payments  $\pi(R)$  ▷ E.g., VCG (Appendix A)
20 return  $a^*(R)$  and  $\pi(R)$ 

```

prices, then the corresponding allocation is efficient and is returned, along with payments $\pi(R)$ according to the deployed payment rule (Line 15). If, by the end of the ML-powered rounds, the market has not cleared, we optionally allow bidders to submit push bids, analogously to the supplementary round of the CCA (Line 17) and calculate the optimal allocation $a^*(R)$ and the payments $\pi(R)$ (Lines 18 and 19). Note that ML-CCA can be combined with various possible payment rules $\pi(R)$, such as VCG or VCG-nearest.

6 Experiments

In this section, we experimentally evaluate the performance of our proposed ML-CCA from Algorithm 2.

6.1 Experiment Setup.

To generate synthetic CA instances, we use the GSVM, LSVM, SRVM, and MRVM domains from the spectrum auction test suite (SATS) (Weiss, Lubin, and Seuken 2017) (see Appendix D.1 for details). We compare our ML-CCA with the original CCA. For both mechanisms, we allow a maximum of 100 clock rounds per instance, i.e., we set $Q^{\text{max}} = 100$. For CCA, we set the price increment to 5% as in (Industry Canada 2013) and optimized the initial reserve

prices to maximize its efficiency. For ML-CCA, we create Q^{init} price vectors to generate the initial demand query data using the same price update rule as the CCA, with the price increment adjusted to accommodate for the reduced number of rounds following this price update rule. In GSVM, LSVM and SRVM we set $Q^{\text{init}} = 20$ for ML-CCA, while in MRVM we set $Q^{\text{init}} = 50$. After each clock round, we report efficiency according to the clock bids up to that round, as well as efficiency if those clock bids were raised (see Section 2.2). Finally, we report the efficiency if the last clock round was supplemented with $Q^{\text{P-Max}} = 100$ bids using the profit max heuristic. Note that this is a very unrealistic and cognitively expensive bidding heuristic in practice, as it requires the agents to both discover their top 100 most profitable bundles as well as report their exact values for them, and thus only adds theoretical value to gauge the difficulty of each domain.

6.2 Hyperparameter Optimization (HPO).

We optimized the hyperparameters (HPs) of the mMVNNs for each bidder type of each domain. Specifically, for each bidder type we trained an mMVNN on the demand responses of a bidder of that type on 50 CCA clock rounds and selected the HPs that resulted in the highest R^2 on validation set 2 as described in Section 3.1. For more details please see Appendix D.3.

6.3 Results.

All efficiency results are presented in Table 1, while in Figure 2 we present the efficiency after each clock round, as well as the efficiency if those clock bids were enhanced with the clock bids raised heuristic (for 95% CIs and p -values see Appendix D.7).

In GSVM, ML-CCA's clock phase exhibits over 7.8% points higher efficiency compared to the CCA, while if we add the clock bids raised heuristic to both mechanisms, ML-CCA still exhibits over 5.3% points higher efficiency. At the same time, ML-CCA is able to find clearing prices in 56% of the instances, as opposed to only 3% for the CCA.

The results for LSVM are qualitatively very similar; ML-CCA's clock phase increases efficiency compared to the CCA by over 9% points, while clearing the market in 26% of the cases as opposed to 0%. If we add the clock bids raised heuristic to both mechanisms, ML-CCA still increases efficiency by over 4.7% points.

The SRVM domain, as suggested by the existence of only 3 unique goods, is quite easy to solve. Thus, both mechanisms can achieve almost 100% efficiency after their clock

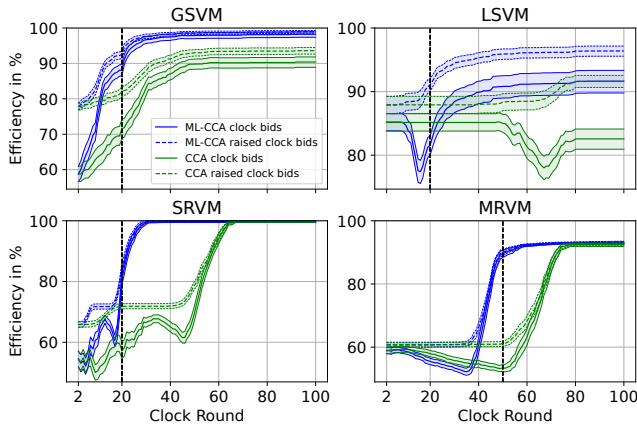


Figure 2: Efficiency path plots in SATS for ML-CCA and CCA both after clock bids (solid lines) and raised clock bids (dashed lines). Averaged over 100 runs including a 95% CI. The dashed black vertical line indicates the value of Q^{init} .

phase. For the clock bids raised heuristic our method reduces the efficiency loss by a factor of more than two (from 0.19% to 0.07%), and additionally, one can see from Figure 2 that our method reaches over 99% in less than 30 rounds.

In MRVM, ML-CCA again achieves statistically significantly better results for all 3 bidding heuristics. Notably, the CCA needs both the clock bids raised heuristic *and* 38 profit max bids to reach the same efficiency as our ML-CCA clock phase, i.e., it needs up to 138 additional *value* queries per bidder (see Appendix D.7). In MRVM, LCPs never exist, thus neither ML-CCA nor the CCA can ever clear the market.

To put our efficiency improvements in perspective, in the GSVM, LSVM and MRVM domains, ML-CCA’s clock phase achieves higher efficiency than the CCA enhanced with the clock bids raised heuristic, i.e., the CCA, even if it uses up to an additional 100 value queries per bidder, cannot match the efficiency of our ML-powered clock phase. In Figure 2, we see that our ML-CCA can (almost) reach the efficiency numbers of Table 1 in a significantly reduced number of clock rounds compared to the CCA, while if we attempt to “speed up” the CCA, then its efficiency can substantially drop, see Appendix D.8. In particular, in GSVM and LSVM, using 50 clock rounds, our ML-CCA can achieve higher efficiency than the CCA can in 100 clock rounds.

6.4 Computational Efficiency.

For our choice of hyperparameters, the computation time when using 8 CPU cores⁵ (see Appendix D.2 for details on our compute infrastructure) for a single round of the ML-CCA averages under 45 minutes for all domains tested. Notably, for three out of four domains, it averages less than 10 minutes (see Table 7 in Appendix D.7). The overwhelming majority of this time is devoted to training the mMVNNs of

⁵Note that Algorithm 1 is not GPU-implementable, as it requires solving a MIP in each iteration, for every demand response by an agent

all bidders using our Algorithm 1 and generating the next DQ using our Algorithm 3 detailed in Section 4. It is important to note that both of these algorithms can always be parallelized by up to the number of bidders N , further reducing the time required. In our implementation, while we parallelized the training of the N mMVNNs, we did not do so for the generation of the next DQ. In spectrum auctions, typically no more than 2 rounds are conducted per day. For the results presented in this paper, we ran the full auction for 400 instances. Given the estimated welfare improvements of over 25 million USD per auction, attributed to ML-CCA’s efficiency gains, we consider ML-CCA’s computational and time requirements to be negligible.

7 Conclusion

We have proposed a novel method for training MVNNs to approximate the bidders’ value functions based on demand query observations. Additionally, we have framed the task of determining the price vector with the highest clearing potential as minimization of an objective function that we prove is convex, Lipschitz-continuous, a.e. differentiable, and whose gradient for linear prices has an intuitive economic interpretation: change the price of every good proportionally to its predicted under/over-demand at the current prices. The resulting mechanism (ML-CCA) from combining these two components exhibits significantly higher clearing potential than the CCA and can increase efficiency by up to 9% points while at the same time converging in a much smaller number of rounds. Thus, we have designed the first *practical* ML-powered auction that employs the same interaction paradigm as the CCA, i.e., demand queries instead of cognitively too complex value queries, yet is able to significantly outperform the CCA in terms of both efficiency and clearing potential in realistic domains.

Acknowledgments

This paper is part of a project that has received funding from the European Research Council (ERC) under the European Union’s Horizon 2020 research and innovation program (Grant agreement No. 805542).

References

- Ausubel, L. M.; and Baranov, O. 2017. A practical guide to the combinatorial clock auction. *Economic Journal*, 127(605): F334–F350. 1, 3, 11, 12, 24
- Ausubel, L. M.; and Baranov, O. 2019. Iterative Vickrey Pricing in Dynamic Auctions. 12
- Ausubel, L. M.; and Baranov, O. 2020. Revealed Preference and Activity Rules in Dynamic Auctions. *International Economic Review*, 61(2): 471–502. 12
- Ausubel, L. M.; and Baranov, O. V. 2014. Market Design and the Evolution of the Combinatorial Clock Auction. *The American Economic Review*, 104(5): 446–451. 12
- Ausubel, L. M.; Cramton, P.; and Milgrom, P. 2006. The clock-proxy auction: A practical combinatorial auction design. In Cramton, P.; Shoham, Y.; and Steinberg, R., eds., *Combinatorial Auctions*, 115–138. MIT Press. 1, 3

- Balcan, M.-F.; Sandholm, T.; and Vitercik, E. 2023. Generalization Guarantees for Multi-item Profit Maximization: Pricing, Auctions, and Randomized Mechanisms. *arXiv:1705.00243*. 2
- Bertsekas, D. P. 1971. *Control of uncertain systems with a set-membership description of the uncertainty*. Ph.D. thesis, Massachusetts Institute of Technology. 17
- Bertsekas, D. P. 1999. *Nonlinear Programming*. Athena scientific optimization and computation series. Athena Scientific. ISBN 9781886529007. 17
- Bichler, M.; Hao, Z.; and Adomavicius, G. 2017. *Coalition-based pricing in ascending combinatorial auctions*, 493–528. Cambridge University Press. ISBN 9781107135345. 1
- Bikhchandani, S.; and Ostroy, J. M. 2002. The package assignment model. *Journal of Economic theory*, 107(2): 377–406. 5, 14
- Blum, A.; Jackson, J.; Sandholm, T.; and Zinkevich, M. 2004. Preference elicitation and query learning. *Journal of Machine Learning Research*, 5: 649–667. 1
- Brero, G.; and Lahaie, S. 2018. A Bayesian clearing mechanism for combinatorial auctions. In *Proceedings of the 32nd AAAI Conference on Artificial Intelligence*. 2
- Brero, G.; Lahaie, S.; and Seuken, S. 2019. Fast Iterative Combinatorial Auctions via Bayesian Learning. In *Proceedings of the 33rd AAAI Conference of Artificial Intelligence*. 2, 3
- Brero, G.; Lubin, B.; and Seuken, S. 2018. Combinatorial Auctions via Machine Learning-based Preference Elicitation. In *Proceedings of the 27th International Joint Conference on Artificial Intelligence*. 1
- Brero, G.; Lubin, B.; and Seuken, S. 2021. Machine Learning-powered Iterative Combinatorial Auctions. *arXiv preprint arXiv:1911.08042*. 1, 3
- Cole, R.; and Roughgarden, T. 2014. The Sample Complexity of Revenue Maximization. In *Proceedings of the Forty-Sixth Annual ACM Symposium on Theory of Computing*, STOC '14, 243–252. New York, NY, USA: Association for Computing Machinery. ISBN 9781450327107. 2
- Cramton, P. 2013. Spectrum auction design. *Review of Industrial Organization*, 42(2): 161–190. 1, 2, 3
- Danskin, J. M. 1967. *The theory of Max-Min and its application to weapons allocation problems [by] John M. Danskin*. Springer-Verlag Berlin, New York. 17
- Dütting, P.; Feng, Z.; Narasimhan, H.; Parkes, D. C.; and Ravindranath, S. S. 2019. Optimal auctions through deep learning. In *Proceedings of the 36th International Conference on Machine Learning*. 2
- Dütting, P.; Fischer, F.; Jirapinyo, P.; Lai, J. K.; Lubin, B.; and Parkes, D. C. 2015. Payment rules through discriminant-based classifiers. *ACM Transactions on Economics and Computation*, 3(1): 5. 2
- Goeree, J. K.; and Holt, C. A. 2010. Hierarchical package bidding: A paper & pencil combinatorial auction. *Games and Economic Behavior*, 70(1): 146–169. 19
- Goetzendorff, A.; Bichler, M.; Shabalin, P.; and Day, R. W. 2015. Compact Bid Languages and Core Pricing in Large Multi-item Auctions. *Management Science*, 61(7): 1684–1703. 1
- Golowich, N.; Narasimhan, H.; and Parkes, D. C. 2018. Deep Learning for Multi-Facility Location Mechanism Design. In *Proceedings of the Twenty-seventh International Joint Conference on Artificial Intelligence and the Twenty-third European Conference on Artificial Intelligence*, 261–267. 2
- Heiss, J. M.; Weisstener, J.; Wutte, H. S.; Seuken, S.; and Teichmann, J. 2022. NOMU: Neural Optimization-based Model Uncertainty. In *Proceedings of the 39th International Conference on Machine Learning*, volume 162 of *Proceedings of Machine Learning Research*, 8708–8758. PMLR. 2
- Industry Canada. 2013. "Responses to Clarification Questions on the Licensing Framework for Mobile Broadband Services (MBS) — 700 MHz Band". 3, 7
- Kwasnica, A. M.; Ledyard, J. O.; Porter, D.; and DeMartini, C. 2005. A New and Improved Design for Multiobject Iterative Auctions. *Management Science*, 51(3): 419–434. 1
- Lahaie, S.; and Lubin, B. 2019. Adaptive-Price Combinatorial Auctions. In *Proceedings of the 2019 ACM Conference on Economics and Computation*, EC '19, 749–750. New York, NY, USA: Association for Computing Machinery. ISBN 9781450367929. 2
- Lahaie, S. M.; and Parkes, D. C. 2004. Applying learning algorithms to preference elicitation. In *Proceedings of the 5th ACM Conference on Electronic Commerce*. 1
- Milgrom, P.; and Segal, I. 2017. Designing the US incentive auction. *Handbook of spectrum auction design*, 803–812. 1
- Morgenstern, J.; and Roughgarden, T. 2015. The Pseudo-Dimension of near-Optimal Auctions. In *Proceedings of the 28th International Conference on Neural Information Processing Systems - Volume 1*, NIPS'15, 136–144. Cambridge, MA, USA: MIT Press. 2
- Narasimhan, H.; Agarwal, S. B.; and Parkes, D. C. 2016. Automated mechanism design without money via machine learning. In *Proceedings of the 25th International Joint Conference on Artificial Intelligence*. 2
- Nisan, N.; and Segal, I. 2006. The communication requirements of efficient allocations and supporting prices. *Journal of Economic Theory*, 129(1): 192–224. 1
- Parkes, D. C.; and Ungar, L. H. 2000. Iterative combinatorial auctions: Theory and practice. In *Proceedings of the seventeenth national Conference on artificial intelligence*, 74–81. 3
- Pogančić, M. V.; Paulus, A.; Musil, V.; Martius, G.; and Ro-linek, M. 2020. Differentiation of blackbox combinatorial solvers. In *International Conference on Learning Representations*. 6
- Rassenti, S. J.; Smith, V. L.; and Bulfin, R. L. 1982. A combinatorial auction mechanism for airport time slot allocation. *The Bell Journal of Economics*, 402–417. 1
- Rockafellar, R. T. 1970. *Convex Analysis*. Princeton: Princeton University Press. ISBN 9781400873173. 17, 18

- Scheffel, T.; Ziegler, G.; and Bichler, M. 2012. On the impact of package selection in combinatorial auctions: an experimental study in the context of spectrum auction design. *Experimental Economics*, 15(4): 667–692. [19](#)
- Soumalias, E.; Weissteiner, J.; Heiss, J.; and Seuken, S. 2024. Machine Learning-powered Combinatorial Clock Auction. *Proceedings of the AAAI Conference on Artificial Intelligence*, 38. [1](#)
- Soumalias, E.; Zamanlooy, B.; Weissteiner, J.; and Seuken, S. 2023. Machine Learning-powered Course Allocation. *arXiv preprint arXiv:2210.00954*. [3](#)
- Weiss, M.; Lubin, B.; and Seuken, S. 2017. Sats: A universal spectrum auction test suite. In *Proceedings of the 16th Conference on Autonomous Agents and MultiAgent Systems*, 51–59. [7](#), [19](#)
- Weissteiner, J.; Heiss, J.; Siems, J.; and Seuken, S. 2022a. Monotone-Value Neural Networks: Exploiting Preference Monotonicity in Combinatorial Assignment. In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, 541–548. International Joint Conferences on Artificial Intelligence Organization. Main Track. [2](#), [4](#), [12](#), [13](#), [14](#), [20](#), [23](#)
- Weissteiner, J.; Heiss, J.; Siems, J.; and Seuken, S. 2023. Bayesian Optimization-based Combinatorial Assignment. *Proceedings of the AAAI Conference on Artificial Intelligence*, 37. [2](#), [4](#), [11](#), [12](#), [14](#), [20](#), [23](#)
- Weissteiner, J.; and Seuken, S. 2020. Deep Learning—Powered Iterative Combinatorial Auctions. *Proceedings of the AAAI Conference on Artificial Intelligence*, 34(02): 2284–2293. [1](#), [12](#), [20](#)
- Weissteiner, J.; Wendler, C.; Seuken, S.; Lubin, B.; and Püschel, M. 2022b. Fourier Analysis-based Iterative Combinatorial Auctions. In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, 549–556. International Joint Conferences on Artificial Intelligence Organization. Main Track. [1](#), [12](#)

Appendix

A Payment and Activity Rules

In this section, we reprint the VCG and VCG-nearest payment rules, as well as give an overview of activity rules for the CCA, and argue why the most prominent choices are also applicable to our ML-CCA.

A.1 VCG Payments from Demand Query Data

Definition A.1. (VCG PAYMENTS FROM DEMAND QUERY DATA) Let $R = (R_1, \dots, R_n)$ denote an elicited set of demand query data from each bidder and let $R_{-i} := (R_1, \dots, R_{i-1}, R_{i+1}, \dots, R_n)$. We then calculate the VCG payments $\pi^{\text{VCG}}(R) = (\pi_1^{\text{VCG}}(R), \dots, \pi_n^{\text{VCG}}(R)) \in \mathbb{R}_{\geq 0}^n$ as follows:

$$\pi_i^{\text{VCG}}(R) := \sum_{j \in N \setminus \{i\}} \langle (p_{R_{-i}}^*)_{j}, (a_{R_{-i}}^*)_{j} \rangle - \sum_{j \in N \setminus \{i\}} \langle (p_R^*)_{j}, (a_R^*)_{j} \rangle. \quad (12)$$

$$(13)$$

where $a_{R_{-i}}^*$ is the allocation that maximizes the inferred social welfare (SW) when excluding bidder i , i.e.,

$$a_{R_{-i}}^* \in \underset{\mathcal{F} \ni (x_j^*(p^{(r_j)}))_{j \in N \setminus \{i\}}}{\operatorname{argmax}} \sum_{j \in N \setminus \{i\}} \langle p^{(r_j)}, x_j^*(p^{(r_j)}) \rangle, \quad (14)$$

$$: r_j \in \{1, \dots, K\} \forall j \in N \setminus \{i\}$$

and $p_{R_{-i}}^* = ((p_{R_{-i}}^*)_1, \dots, (p_{R_{-i}}^*)_n) \in \mathbb{R}_{\geq 0}^{mn}$ denote the corresponding price vectors that lead to $a_{R_{-i}}^*$, and a_R^* is a inferred-social-welfare-maximizing allocation (see Equation (2)) with corresponding prices $p_R^* = ((p_R^*)_1, \dots, (p_R^*)_n) \in \mathbb{R}_{\geq 0}^{mn}$.

Thus, when using VCG, bidder i 's utility is:

$$u_i = v_i((a_R^*)_i) - \pi_i^{\text{VCG}}(R) = v_i((a_R^*)_i) + \sum_{j \in N \setminus \{i\}} \langle (p_R^*)_{j}, (a_R^*)_{j} \rangle - \sum_{j \in N \setminus \{i\}} \langle (p_{R_{-i}}^*)_{j}, (a_{R_{-i}}^*)_{j} \rangle.$$

Remark A.1. Note that Definition A.1 defines VCG payments for a set of elicited demand query data R using the bidders' inferred values from this data set. Specifically, those would be the VCG payments after the CCA's clock phase for example (i.e., if there was no supplementary round). However, one can analogously define VCG-payments for any given set of reported bundle-value pairs (see Weissteiner et al. (2023, Definition B.1.)). For example, in the case of additional value bids, such as supplementary round push bids, one would use Weissteiner et al. (2023, Definition B.1.) and set $\hat{v}_i(x^r) = \max\{x^r, p^r\}, b_i^s(x^r)\}$, where b_i^s is the bidder's supplementary round bid for that bundle (or zero, if she did not bid on it in the supplementary round), i.e., bidder i 's bid for any bundle is the maximum of her largest inferred value for that bundle based on the clock round bids and the supplementary round bids.

A.2 VCG-Nearest Payments

To define the VCG-nearest payments, we must first introduce the core:

Definition A.2. (THE CORE) An outcome $(a, \pi) \in \mathcal{F} \times \mathbb{R}_{\geq 0}^n$ (i.e., a tuple of a feasible allocation a and payments π) is in the core if it satisfies the following two properties:

1. The outcome is individual rational, i.e., $u_i = v_i(a_i) - \pi_i \geq 0$ for all $i \in N$
2. The core constraints

$$\forall L \subseteq N \quad \sum_{i \in N \setminus L} \pi_i(R) \geq \max_{a' \in \mathcal{F}} \sum_{i \in L} v_i(a'_i) - \sum_{i \in L} v_i(a_i) \quad (15)$$

where $v_i(a_i)$ is bidder i 's value for bundle a_i and $\mathcal{F}$ is the set of feasible allocations.

In words, a payment vector π (together with a feasible allocation a) is in the core if no coalition of bidders $L \subset N$ is willing to pay more for the items than the mechanism is charging the winners. Note that by replacing the true values $v_i(a_i)$ with the bidders' (possibly untruthful) bids $b_i(a_i)$ in Definition A.2 one can equivalently define the revealed core.

Now, we can define

Definition A.3. (MINIMUM REVENUE CORE) Among all payment vectors in the (revealed) core, the (revealed) minimum revenue core is the set of payment vectors with smallest L_1 -norm, i.e., which minimize the sum of the payments of all bidders.

We can now define VCG-nearest payments:

Definition A.4. (VCG-NEAREST PAYMENTS) Given an allocation a_R for bidder reports R , the VCG-nearest payments $\pi^{\text{VCG-nearest}}(R)$ are defined as the vector of payments in the (revealed) minimum revenue core that minimizes the L_2 -norm to the VCG payment vector $\pi^{\text{VCG}}(R)$.

A.3 On the Importance of Activity Rules to Align Incentives

In the CCA, activity rules serve multiple purposes. First, they can help speed up the auction process. Second, they reduce "bid-sniping" opportunities, i.e., bidders concealing their true intentions until the very last rounds of the auction.⁶ Third, they can limit surprise bids in the supplementary round of the CCA and significantly reduce a bidder's ability to drive up her opponents payments by overbidding on bundles that she can no longer win (Ausubel and Baranov 2017). There are two types of activity rules that are implemented in a CCA:

1. Clock phase activity rules, that limit the bundles that an agent can bid on during the clock phase, based on her bids in previous clock rounds.
2. Supplementary round activity rules, that restrict the amount that an agent can bid on for various sets of items during the supplementary round.

⁶The notion of "bid-sniping" first originated in eBay auctions with predetermined ending times, where the high-value bidder can sometimes reduce her payments by submitting her bid at the very last moment.

Most of the activity rules that were traditionally used for the clock phase of the CCA were based on either revealed-preference considerations or some *points-based system*, where the main idea is to assign points to each item prior to the auction, and only allow bidders to submit monotonically non-increasing in points bids, i.e., as the rounds progress and the prices increase, the bidders cannot submit bids for larger sets of items. Both of these approaches, as well as hybrid combinations thereof, were shown to actually further interfere with truthful bidding in some cases (Ausubel and Baranov 2014, 2020).

However, Ausubel and Baranov (2019) showed that basing the clock phase activity rule not on the above but instead entirely upon the *generalized axiom of revealed preference (GARP)* can dynamically approximate VCG payoffs and thus improve the bidding incentives of the CCA. GARP imposes revealed-preference constraints (see Definition A.5) to the bidder’s demand responses, i.e., the GARP activity rule requires the bidder to exhibit rational behaviour in her demand choices. Importantly, the GARP activity rule *does not* require a monotonic price trajectory. Thus, it can also be applied in our ML-powered clock phase, allowing the clock phase of our ML-CCA to enjoy the same improvement in bidding incentives.

For the supplementary round, the CCA’s most prominent activity rules are again based on a combination of the same points-based system and revealed-preference ideas. For this, we need to define the following constraint:

Definition A.5. (REVEALED-PREFERENCE CONSTRAINT)
The revealed-preference constraint for bundle $x \in X$ with respect to clock round r is

$$b_i(x) \leq b_i(x^r) + \langle p^r, x - x^r \rangle, \quad (16)$$

where $b_i(x) \in \mathbb{R}_{\geq 0}$ is bidder i ’s bid for bundle $x \in X$ in the supplementary round, $x^r \in \mathcal{X}$ is the bundle demanded by the agent at clock round r , $b_i(x^r) \in \mathbb{R}_{\geq 0}$ is the final bid for bundle $x^r \in \mathcal{X}$ and $p^r \in \mathbb{R}_{\geq 0}^m$ is the linear price vector of clock round r .

Intuitively, the revealed-preference constraint states that a bidder is not allowed to claim a high value for bundle $x \in \mathcal{X}$ relative to bundle $x^r \in \mathcal{X}$, given that she claimed to prefer bundle $x^r \in \mathcal{X}$ at clock round r (see Inequality (6)). The difference between the three most prominent supplementary round activity rules is with respect to *which* clock rounds the revealed-preference constraint should be satisfied. Specifically:

1. *Final Cap*: A bid for bundle $x \in \mathcal{X}$ should satisfy the revealed-preference constraint (Definition A.5) with respect to the *final* clock round’s price $p^{Q^{\max}} \in \mathbb{R}_{\geq 0}$ and bundle $x^{Q^{\max}} \in \mathcal{X}$.
2. *Relative Cap*: A bid for bundle $x \in \mathcal{X}$ should satisfy the revealed-preference constraint (Definition A.5) with respect to the last clock round for which the bidder was eligible for that bundle $x \in \mathcal{X}$, based on the points-based system.
3. *Intermediate Cap*: A bid for bundle $x \in \mathcal{X}$ should satisfy the revealed-preference constraint (Definition A.5) with

respect to all eligibility-reducing rounds, starting from the last clock round for which the bidder was eligible for $x \in \mathcal{X}$ based on the point system.

Ausubel and Baranov (2017) showed that combining the *Final Cap* and *Relative Cap* activity rules leads to the largest amount of reduction in bid-sniping opportunities for the UK 4G auction, as measured by the theoretical bid amount that each bidder would need to increase her bid by in the supplementary round in order to protect her final clock round bundle. Finally, note that the *Final*- and *Intermediate Cap* activity rules can also be applied to our ML-CCA.⁷

To conclude, we observe that for both its clock phase and the supplementary round, our ML-CCA (with the same F_{init} method as the CCA), is compatible with the most prominent activity rules for the corresponding phases of the CCA, while it is also obviously compatible with the most prominent payment rule, VCG-nearest prices (Definition A.4). This, combined with the fact that the ML-CCA has the same interaction paradigm for the bidders as the CCA, is a very strong indication that our ML-CCA can reduce the opportunities for the bidders to misreport to a similar extent as the classical CCA.

B Multiset MVNNs

B.1 Advantages of mMVNNs over MVNNs

In Weissteiner and Seuken (2020) and Weissteiner et al. (2022b,a, 2023) items with a capacity c_k , were treated as c_k distinct items, without exploiting the prior knowledge that these c_k items are indistinguishable to the bidders. Sufficiently large classical MVNNs that were trained on large enough training sets would at some point learn that these items are indistinguishable, but this prior information was not incorporated at an architectural level. For multiset MVNNs (mMVNNs) this prior knowledge is hard-coded directly into the architecture (see Definition 1). This additional prior knowledge is particularly beneficial for small training data sets in terms of generalization. At the same time, it can significantly reduce the dimensionality of the input space, as Example B.1 illustrates.

Example B.1. If we have 30 different items each of capacity 2 (i.e., $\mathcal{X} = \{0, 1, 2\}^{30}$), then the bundle $x = (1, \dots, 1) \in \mathcal{X}$ containing one of each 30 items would have $2^{30} > 1$ billion different representations as sets corresponding to binary vectors in $\tilde{\mathcal{X}} = \{0, 1\}^{60}$ (where the items are treated as 60 items of capacity 1). All these 2^{30} representations in $\tilde{\mathcal{X}}$ could be mapped to different values by a classical MVNN $\tilde{\mathcal{M}}_i^\theta : \tilde{\mathcal{X}} \rightarrow \mathbb{R}_{\geq 0}$, while we actually have the hard prior knowledge that they should all have exactly the same value due to indistinguishable items. And all these 2^{30} representations in $\tilde{\mathcal{X}}$ correspond to the same multiset

⁷With the modification for the *Relative Cap* rule that the revealed-preference constraint should hold for the Q^{init} rounds that follow the same price update rule as the CCA, and then the ML-powered clock rounds should be treated as corresponding to the same amount of points, since the prices in these rounds on aggregate stay very close to the prices of the last Q^{init} round, as shown in Figures 5 to 8.

$x = (1, \dots, 1) \in \mathcal{X}$, which gets assigned to exactly one value $\mathcal{M}_i^\theta(x)$ by an mMVNN $\mathcal{M}_i^\theta : \mathcal{X} \rightarrow \mathbb{R}_{\geq 0}$.

Furthermore, the solution times of MILPs are likely to benefit from the mMVNNs' reduced variable count compared to traditional MVNNs. In the case of classical MVNNs, it was necessary to introduce one binary variable for each indistinguishable duplicate of an item, resulting in $\tilde{m} = \sum_{k=1}^m c_k$. In contrast, mMVNNs require only m integer variables. For an experimental comparison of mMVNNs to MVNNs, see Appendix D.10.

B.2 Universality of mMVNNs

Note that Weissteiner et al. (2022a) have proven universality of MVNNs only for the case of binary input vectors corresponding to classical sets (i.e., $c_k = 1 \forall k \in M$). Here, we prove in Appendix B.2 universality of mMVNNs for arbitrary capacities $c \in \mathbb{N}^m$ corresponding to multiset domains such as our $\mathcal{X}$.

First, we recall the following definition:

Definition B.1. The set $\mathcal{V}$ of all monotonic⁸ and normalized functions from $\mathcal{X}$ to $\mathbb{R}_{\geq 0}$ is defined as

$$\mathcal{V} := \{v_i : \mathcal{X} \rightarrow \mathbb{R}_+ \mid \text{satisfy (N) and (M)}\}. \quad (17)$$

The following lemma says that every mMVNN is monotonic (M) and normalized (N) in the sense of Definition B.1.

Lemma 1. Let $\mathcal{M}_i^\theta : \mathcal{X} \rightarrow \mathbb{R}_+$ be an mMVNN from Definition 1. Then it holds that $\mathcal{M}_i^{(W^i, b^i)} \in \mathcal{V}$ for all $W^i \geq 0$ and $b^i \leq 0$.

Proof. The proof of this lemma is perfectly analogous to the proof of (Weissteiner et al. 2022a, Lemma 1). $\square$

Now we are ready to present a constructive proof for Theorem 1.

Proof of Theorem 1

Proof. This proof follows a similar strategy as the proof of (Weissteiner et al. 2022a, Theorem 1).

$$1. \mathcal{V} \supseteq \left\{ \mathcal{M}_i^{(W^i, b^i)} : W^{i,k} \geq 0, b^{i,k} \leq 0 \forall k \in \{1, \dots, K_i\} \right\}$$

This direction follows immediately from Lemma 1.

$$2. \mathcal{V} \subseteq \left\{ \mathcal{M}_i^{(W^i, b^i)} : W^{i,k} \geq 0, b^{i,k} \leq 0 \forall k \in \{1, \dots, K_i\} \right\}$$

Let $(v_i(x))_{x \in \mathcal{X}} \in \mathcal{V}$. For the reverse direction, we give a constructive proof, i.e., we construct an mMVNN $\mathcal{M}_i^\theta$ with $\theta = (W_{v_i}^i, b_{v_i}^i)_{i \in \{1, \dots, 4\}}$ such that $\mathcal{M}_i^\theta(x) = v_i(x)$ for all $x \in \mathcal{X}$.

⁸Within this paper, by “monotonic”, we always refer to weakly monotonically increasing (M), i.e., monotonically non-decreasing. In multiset-notation this reads as: $\forall a, b \in \mathcal{X} : (a \subseteq b \implies v_i(a) \leq v_i(b))$.

Let $(w_j)_{j=1}^{|\mathcal{X}|}$ denote the values corresponding to $(v_i(x))_{x \in \mathcal{X}}$ of all $|\mathcal{X}| = \prod_{j=1}^m (c_j + 1)$ possible bundles $x \in \mathcal{X}$ sorted by value in increasing order, i.e., let $x_1 = (0, \dots, 0)$ with

$$w_1 := v_i(x_1) = 0, \quad (18)$$

let $x_{|\mathcal{X}|} = c = (c_1, \dots, c_m)$ with

$$w_{|\mathcal{X}|} := v_i(x_{|\mathcal{X}|}), \quad (19)$$

and $x_j, x_l \in \mathcal{X} \setminus \{x_1, x_{|\mathcal{X}|}\}$ for $1 < l \leq j \leq |\mathcal{X}| - 1$ with

$$w_j := v_i(x_j) \leq w_l := v_i(x_l). \quad (20)$$

In the following, for $x_l, x_j \in \mathcal{X}$, we use the notation $x_l \leq x_j$ iff $x_{l,k} \leq x_{j,k} \forall k \in M$. Thus, we write $x_l \not\leq x_j$ iff $\exists k \in M : x_{l,k} > x_{j,k}$, i.e., iff $x_l \leq x_j$ is not the case.⁹ Furthermore, we denote by $\langle \cdot, \cdot \rangle$ the Euclidean scalar product on $\mathbb{R}^m$. Then, for all $x \in \mathcal{X}$:

$$v_i(\xi) = \sum_{l=1}^{|\mathcal{X}|-1} (w_{l+1} - w_l) \mathbb{1}_{\{\forall j \in \{1, \dots, l\} : \xi \not\leq x_j\}} \quad (21)$$

$$= \sum_{l=1}^{|\mathcal{X}|-1} (w_{l+1} - w_l) \varphi_{0,1} \left(\sum_{j=1}^l \varphi_{0,1} \left(\sum_{k=1}^m \varphi_{0,1}(\xi_k - x_{j,k}) \right) - (l-1) \right), \quad (22)$$

where the second equality follows since

$$\xi \not\leq x_j \iff \exists k \in M : \xi_k > x_{j,k} \quad (23)$$

$$\iff \exists k \in M : \varphi_{0,1}(\xi_k - x_{j,k}) = 1 \quad (24)$$

$$\iff \varphi_{0,1} \left(\sum_{k=1}^m \varphi_{0,1}(\xi_k - x_{j,k}) \right) = 1, \quad (25)$$

which implies that

$$\begin{aligned} &(\forall j \in \{1, \dots, l\} : \xi \not\leq x_j) \\ &\iff \sum_{j=1}^l \varphi_{0,1} \left(\sum_{k=1}^m \varphi_{0,1}(\xi_k - x_{j,k}) \right) = l, \end{aligned} \quad (26)$$

and

$$\begin{aligned} &\mathbb{1}_{\{\forall j \in \{1, \dots, l\} : \xi \not\leq x_j\}} \\ &= \varphi_{0,1} \left(\sum_{j=1}^l \varphi_{0,1} \left(\sum_{k=1}^m \varphi_{0,1}(\xi_k - x_{j,k}) \right) - (l-1) \right). \end{aligned} \quad (27)$$

To match the structure of the mMVNN architecture defined in Definition 1, we can write $v_i(\xi) = v_i(D^{-1}D\xi)$ and plug in $D^{-1}D\xi$ instead of ξ in Equation (22).¹⁰

⁹Note, that in the multidimensional case $m > 1$, the two symbols $\not\leq$ and $>$ have a different meaning, i.e., $x_l > x_j \iff (x_l \not\leq x_j \text{ and } x_l \geq x_j)$. Further note that $<, \leq$ and $\not\leq$ exactly correspond to $\subset, \subseteq$ and $\not\subseteq$ respectively in multiset notation.

¹⁰The diagonal matrix $D := \text{diag}(1/c_1, \dots, 1/c_m)$ is invertible with $D^{-1} = \text{diag}(c_1, \dots, c_m)$, since all capacities c_i are strictly larger than 0.

Equation (22) can be rewritten as the following mMVNN

$$\begin{aligned} \mathcal{M}_i^\theta(x) = & W_{v_i}^{i,4} \varphi_{0,1} \left(W_{v_i}^{i,3} \varphi_{0,1} \left(W_{v_i}^{i,2} \varphi_{0,1} \left(W_{v_i}^{i,1} D\xi + b_{v_i}^{i,1} \right) + b_{v_i}^{i,2} \right) b_{v_i}^{i,3} \right) \\ & + b_{v_i}^{i,4} \end{aligned} \quad (28)$$

in the matrix-notation of Definition 1 with weight matrices and bias-vectors be given as:

$$W_{v_i}^{i,1} := \begin{bmatrix} D^{-1} \\ \vdots \\ D^{-1} \end{bmatrix} \in \mathbb{R}^{m(|\mathcal{X}|-1) \times m}, \quad (29)$$

$$b_{v_i}^{i,1} := - \begin{bmatrix} x_1 \\ x_2 \\ \vdots \\ x_{|\mathcal{X}|-1} \end{bmatrix} \in \mathbb{R}^{m(|\mathcal{X}|-1)}, \quad (30)$$

$$\begin{aligned} W_{v_i}^{i,2} := & \begin{bmatrix} \overbrace{1, \dots, 1}^m & \overbrace{0, \dots, 0}^m & \dots & \overbrace{0, \dots, 0}^m \\ 0, \dots, 0 & 1, \dots, 1 & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0, \dots, 0 \\ 0, \dots, 0 & \dots & 0, \dots, 0 & 1, \dots, 1 \end{bmatrix} \\ & \in \mathbb{R}^{(|\mathcal{X}|-1) \times m(|\mathcal{X}|-1)}, \end{aligned} \quad (31)$$

$$b_{v_i}^{i,2} := 0 \in \mathbb{R}^{|\mathcal{X}|-1}, \quad (32)$$

$$W_{v_i}^{i,3} := \begin{bmatrix} 1 & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ \vdots & & \ddots & 0 \\ 1 & \dots & \dots & 1 \end{bmatrix} \in \mathbb{R}^{(|\mathcal{X}|-1) \times (|\mathcal{X}|-1)}, \quad (33)$$

$$b_{v_i}^{i,3} := \begin{bmatrix} 0 \\ -1 \\ \vdots \\ -(|\mathcal{X}| - 2) \end{bmatrix} \in \mathbb{R}^{|\mathcal{X}|-1}, \quad (34)$$

$$W_{v_i}^{i,4} := \begin{bmatrix} w_2 - w_1 \\ w_3 - w_2 \\ \vdots \\ w_{|\mathcal{X}|} - w_{|\mathcal{X}|-1} \end{bmatrix}^\top \in \mathbb{R}^{1 \times (|\mathcal{X}|-1)}, \quad (35)$$

$$b_{v_i}^{i,4} := 0 \in \mathbb{R}, \quad (36)$$

where $W_{v_i}^{i,2} = (\mathbb{1}_{\{(j-1)m < k \leq jm\}})_{1 \leq j \leq |\mathcal{X}|-1, 1 \leq k \leq m(|\mathcal{X}|-1)}$ is an alternative notation to describe Equation (31). Thus, $\mathcal{M}_i^\theta(x)$ is an mMVNN from Definition 1 with 6 layers in total (i.e., 1 input layer, 1 linear normalization layer D , 3 non-linear hidden layers and 1 output layer) and respective dimensions $[m, m, m(|\mathcal{X}|-1), |\mathcal{X}|-1, |\mathcal{X}|-1, 1]$.

□

Remark B.1 (Normalization Layer). *Note that our proof of Theorem 1 would also work without the linear normalization layer D . The normalization layer has the advantage that we can use similar hyperparameters as for classical MVNNs.*

E.g., we can use the same values of the parameters in the initialization scheme as provided in Weissteiner et al. (2023, Section 3.2 and Appendix E).

Remark B.2 (Number of Hidden Layers). *Note that for binary input vectors (corresponding to classical sets, i.e., $c = (1, \dots, 1)$), 2 non-linear hidden layers were sufficient (Weissteiner et al. 2022a, Proof of Theorem 1), while for integer-valued input vectors (corresponding to multisets with capacities $c \in \mathbb{N}^m$), we used 3 non-linear hidden layers for our proof of Theorem 1. It is an interesting open question if 2 non-linear hidden layers would already be sufficient also for our multiset setting with capacities $c \in \mathbb{N}^m$.¹¹*

Remark B.3 (Number of Hidden Neurons). *The dimension $[m, m, m(|\mathcal{X}|-1), |\mathcal{X}|-1, |\mathcal{X}|-1, 1]$ of the mMVNN used in the proof of Theorem 1 is just an upper bound. This proof does not imply that such large networks are actually needed. From a theoretical perspective it is interesting that this upper bound is finite, while for NNs on continuous domains no finite upper bound for perfect approximation exists. In practice much smaller networks are usually sufficient. All the results reported in this paper were achieved by much smaller networks with at most 30 neurons per layer (see Table 2 in Appendix D.3).*

C ML-powered Demand Query Generation: Theoretical Results

In this section, we prove Theorems 2 and 3.

C.1 Proof of Theorem 2

In this subsection, we first present and prove in Lemma 2 an unconstrained version of Theorem 2, which we later use to prove our main statement from Theorem 2.

Lemma 2 extends (Bikhchandani and Ostroy 2002, Theorem 3.1.). Concretely, we additionally show that if linear clearing prices exist, every minimizer of W is a clearing price (while Bikhchandani and Ostroy (2002, Theorem 3.1.) only showed that every clearing price p minimizes W not specifying if there could be other minimizers of W which are not clearing prices).

Lemma 2. *Consider the notation from Definitions 2 and 3 and the objective function $W(p, v) := R(p) + \sum_{i \in N} U(p, v_i)$. Then it holds that, if a linear clearing price exists, every price vector*

$$p' \in \underset{\tilde{p} \in \mathbb{R}_{\geq 0}^m}{\operatorname{argmin}} W(\tilde{p}, v) \quad (37)$$

is a clearing price and the corresponding allocation $a(p')$ in $\mathcal{F}$ is efficient.

Proof of Lemma 2. We first show in Item 1 that for clearing prices $p \in \mathbb{R}_{\geq 0}^m$, the corresponding allocation $a(p) \in$

¹¹Our proof of Theorem 1 works with 1 input layer, 1 linear normalization layer D (which is optional, see Remark B.1), 3 non-linear hidden layers and 1 output layer (i.e., 5 or 6 layers in total). Whereas the (Weissteiner et al. 2022a, Proof of Theorem 1) only required 1 input layer, 2 non-linear hidden layers and 1 output layer (i.e., 4 layers in total).

$\mathcal{F}$ is *efficient*. Next, in Item 2 we show that every linear clearing price $p \in \mathbb{R}_{\geq 0}^m$ minimizes W , i.e., formally $p \in \operatorname{argmin}_{\tilde{p} \in \mathbb{R}_{\geq 0}^m} W(\tilde{p}, v)$. Finally, in Item 3 we show that if linear clearing prices exist, *any* minimizer of W is in fact a linear clearing price.

1. Let $(p, a(p)) \in \mathbb{R}_{\geq 0}^m \times \mathcal{F}$ be clearing prices and the corresponding supported allocation (see Definition 3). Furthermore, let $\tilde{a} \in \mathcal{F}$ be any other feasible allocation. Then it holds that:

$$\sum_{i=1}^n v_i(a_i(p)) = \quad (38)$$

$$\sum_{i=1}^n [v_i(a_i(p)) - \langle p, a_i(p) \rangle] + \sum_{i=1}^n \langle p, a_i(p) \rangle = \quad (39)$$

$$\left(\sum_{i=1}^n U(p, v_i) \right) + R(p) \geq \quad (40)$$

$$\sum_{i=1}^n [v_i(\tilde{a}_i) - \langle p, \tilde{a}_i \rangle] + \sum_{i=1}^n \langle p, \tilde{a}_i \rangle = \quad (41)$$

$$\sum_{i=1}^n v_i(\tilde{a}_i), \quad (42)$$

where the first inequality follows since $(p, a(p))$ are clearing prices and fulfill Items 1 and 2 in Definition 3 and because $\tilde{a} \in \mathcal{F}$ is another feasible allocation. This shows that a supported allocation $a(p)$ is efficient.

2. Let $(p, a(p)) \in \mathbb{R}_{\geq 0}^m \times \mathcal{F}$ be clearing prices and the corresponding supported allocation (see Definition 3). Furthermore, let $\tilde{p} \in \mathbb{R}_{\geq 0}^m$ be any other linear prices. Then it holds that:

$$W(p, v) = \quad (43)$$

$$\max_{a \in \mathcal{F}} \left\{ \sum_{i=1}^n \langle p, a_i \rangle \right\} + \sum_{i=1}^n \max_{x \in \mathcal{X}} \{v_i(x) - \langle p, x \rangle\} = \quad (44)$$

$$\sum_{i=1}^n \langle p, a_i(p) \rangle + \sum_{i=1}^n [v_i(a_i(p)) - \langle p, a_i(p) \rangle] = \quad (45)$$

$$\sum_{i=1}^n v_i(a_i(p)) = \quad (46)$$

$$\sum_{i=1}^n \langle \tilde{p}, a_i(p) \rangle + \sum_{i=1}^n [v_i(a_i(p)) - \langle \tilde{p}, a_i(p) \rangle] \leq \quad (47)$$

$$\max_{a \in \mathcal{F}} \left\{ \sum_{i=1}^n \langle \tilde{p}, a_i \rangle \right\} + \sum_{i=1}^n \max_{x \in \mathcal{X}} \{v_i(x) - \langle \tilde{p}, x \rangle\} = \quad (48)$$

$$W(\tilde{p}, v), \quad (49)$$

where Equation (44) follows by definition of W and Equation (45) follows since $(p, a(p))$ are clearing prices and the corresponding supported allocation, respectively. Thus, we get that

$$W(p, v) \leq W(\tilde{p}, v), \quad (50)$$

for all linear prices $\tilde{p} \in \mathbb{R}_{\geq 0}^m$, which concludes the proof.

3. Let $p' \in \operatorname{argmin}_{\tilde{p} \in \mathbb{R}_{\geq 0}^m} W(\tilde{p}, v)$ be a minimizer of W . Moreover, let $p \in \mathbb{R}_{\geq 0}^m$ denote a linear clearing price and let $a(p) \in \mathcal{F}$ be the corresponding supported allocation (see Definition 3). Furthermore, let $x_i^*(p') \in \operatorname{argmax}_{x \in \mathcal{X}} \{v_i(x) - \langle p', x \rangle\} \forall i \in N$. We know from Item 2 that $W(p', v) = W(p, v)$. Furthermore, we get that

$$\sum_{i=1}^n v_i(x_i^*(p')) - \sum_{i=1}^n \langle p', x_i^*(p') \rangle + \langle p', c \rangle = \quad (51a)$$

$$W(p', v) = W(p, v) = \quad (51b)$$

$$\sum_{i=1}^n v_i(a_i(p)) - \sum_{i=1}^n \langle p, a_i(p) \rangle + \langle p, c \rangle \stackrel{(45)}{=} \quad (51c)$$

$$\sum_{i=1}^n v_i(a_i(p)) = \quad (51d)$$

$$\sum_{i=1}^n v_i(a_i(p)) - \sum_{i=1}^n \langle p', a_i(p) \rangle + \langle p', c \rangle, \quad (51e)$$

where the last equation follows since $a(p)$ is a clearing allocation supported by p and thus $\sum_{i=1}^n \langle p', a_i(p) \rangle = \langle p', c \rangle$. Next, we can subtract $\langle p', c \rangle$ from both sides of Equation (51) to obtain

$$\begin{aligned} & \sum_{i=1}^n v_i(x_i^*(p')) - \langle p', x_i^*(p') \rangle = \\ & \sum_{i=1}^n v_i(a_i(p)) - \langle p', a_i(p) \rangle. \end{aligned} \quad (52)$$

Moreover, from the optimality of $x_i^*(p')$ it follows that

$$v_i(x_i^*(p')) - \langle p', x_i^*(p') \rangle \geq v_i(a_i(p)) - \langle p', a_i(p) \rangle, \quad (53)$$

holds for every $i \in \{1, \dots, n\}$.

Using now Equations (52) and (53), it follows that for every $i \in \{1, \dots, n\}$

$$U(p', v_i) = v_i(x_i^*(p')) - \langle p', x_i^*(p') \rangle = \quad (54)$$

$$v_i(a_i(p)) - \langle p', a_i(p) \rangle. \quad (55)$$

Taken all together, since $\sum_{i \in N} \langle p', a_i(p) \rangle = \sum_{j \in M} \langle p'_j, c_j \rangle = R(p')$ and $U(p', v_i) = v_i(a_i(p)) - \langle p', a_i(p) \rangle$ for all $i \in N$, we can conclude that $(p', a(p)) \in \mathbb{R}_{\geq 0}^m \times \mathcal{F}$ by Definition 3 is a linear clearing price with corresponding supported allocation $a(p) \in \mathcal{F}$. This finalizes the proof of Lemma 2. $\square$

Now, we are ready to prove Theorem 2, which follows almost immediately from Lemma 2.

Proof of Theorem 2. Let p' be a solution to the constrained minimization problem defined by Equation (9). Moreover,

let $p \in \mathbb{R}_{\geq 0}^m$ denote a linear clearing price vector and let $a(p) \in \mathcal{F}$ be the corresponding supported allocation (see Definition 3).

In Item 2 in the proof of Lemma 2, we have seen that p minimizes W without any constraints. The clearing price vector p obviously satisfies constraint (9b), thus p is also a solution to the constrained optimization problem (9). Therefore, in the case that linear clearing prices exist the minimal objective value of the constrained optimization problem from (9) is equal to the minimal objective value of the unconstrained minimization problem (37). From this we can conclude that in the case that linear clearing prices exist, every solution p' of (9) is also a solution of the optimization problem (37). Finally, Lemma 2 tells us that every solution p' of the optimization problem (37) is a clearing price and $a(p')$ is efficient. $\square$

In Theorem 2 we proved that the constraint minimizer of W has the same economical favourable properties as the unconstrained minimizer of W if linear clearing prices exist. Thus, we do not lose anything in cases where linear clearing prices exist, while due to the constraint (9b) we have the advantage of receiving feasible allocations in the case that no linear clearing price exist.

C.2 Details: Proof of Theorem 3

Before we start with the proof, we will quickly provide a proof for Equation (8), since this equation (that we have not explicitly proven in the main paper) is essential in the proof of Theorem 3.

Proof of Equation (8).

$$R(p) := \max_{a \in \mathcal{F}} \left\{ \sum_{i \in N} \langle p, a_i \rangle \right\} \quad (56)$$

$$= \max_{a \in \mathcal{F}} \left\{ \sum_{i \in N} \sum_{j \in M} p_j a_{i,j} \right\} \quad (57)$$

$$= \max_{a \in \mathcal{F}} \left\{ \sum_{j \in M} \sum_{i \in N} p_j a_{i,j} \right\} \quad (58)$$

$$= \max_{a \in \mathcal{F}} \left\{ \sum_{j \in M} p_j \sum_{i \in N} a_{i,j} \right\} \quad (59)$$

$$= \max \left\{ \sum_{j \in M} p_j \sum_{i \in N} a_{i,j} : \sum_{i \in N} a_{i,j} \leq c_j \ \forall j \in M \right\} \quad (60)$$

$$= \sum_{j \in M} p_j c_j = \sum_{j \in M} c_j p_j = \langle c, p \rangle. \quad (61)$$

$\square$

In the first lemma we prove that $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ is Lipschitz-continuous. While neither $p \mapsto x_i^*(p)$, nor $p \mapsto \mathcal{M}_i^\theta(x_i^*(p))$ nor $p \mapsto \langle p, x_i^*(p) \rangle$ are continuous, (surprisingly) $p \mapsto \mathcal{M}_i^\theta(x_i^*(p)) + \langle p, x_i^*(p) \rangle$ is continuous.

Lemma 3 (Continuity). *The map $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ from $\mathbb{R}_{\geq 0}^m$ to $\mathbb{R}_{\geq 0}$ is Lipschitz-continuous with Lipschitz-constant $(n+1)\|c\|_2$.¹²*

Proof. Since $W(p, (\mathcal{M}_i^\theta)_{i=1}^n) := R(p) + \sum_{i \in N} U(p, \mathcal{M}_i^\theta)$ is the sum of $1+n$ functions, we first quickly show that R is Lipschitz-continuous in p and afterwards we will show that also U is Lipschitz-continuous in p .

From Equation (8) it follows that $R(p) = \sum_{j \in M} c_j p_j$ is linear in p and thus Lipschitz-continuous with Lipschitz-constant $\|c\|_2$.

In the remainder of the proof we are going to show that $U(p, \mathcal{M}_i^\theta) := \max_{x \in \mathcal{X}} \{ \mathcal{M}_i^\theta(x) - \langle p, x \rangle \}$ is Lipschitz-continuous in p . Let $p, \tilde{p} \in \mathbb{R}_{\geq 0}^m$ be two price vectors, then we have

$$U(\tilde{p}, \mathcal{M}_i^\theta) \geq \mathcal{M}_i^\theta(\hat{x}_i^*(p)) + \langle \tilde{p}, \hat{x}_i^*(p) \rangle \quad (62)$$

$$= \mathcal{M}_i^\theta(\hat{x}_i^*(p)) + \langle \tilde{p} + p - p, \hat{x}_i^*(p) \rangle \quad (63)$$

$$= \mathcal{M}_i^\theta(\hat{x}_i^*(p)) + \langle p, \hat{x}_i^*(p) \rangle + \langle \tilde{p} - p, \hat{x}_i^*(p) \rangle \quad (64)$$

$$= U(p, \mathcal{M}_i^\theta) + \langle \tilde{p} - p, \hat{x}_i^*(p) \rangle \quad (65)$$

$$\geq U(p, \mathcal{M}_i^\theta) - \|\tilde{p} - p\|_2 \|\hat{x}_i^*(p)\|_2 \quad (66)$$

$$\geq U(p, \mathcal{M}_i^\theta) - \|\tilde{p} - p\|_2 \|c\|_2, \quad (67)$$

and thus

$$U(\tilde{p}, \mathcal{M}_i^\theta) - U(p, \mathcal{M}_i^\theta) \geq -\|\tilde{p} - p\|_2 \|c\|_2. \quad (68)$$

By exchanging the roles of p and $\tilde{p}$, we also obtain $U(p, \mathcal{M}_i^\theta) \geq U(\tilde{p}, \mathcal{M}_i^\theta) - \|\tilde{p} - p\|_2 \|c\|_2$ and thus,

$$U(\tilde{p}, \mathcal{M}_i^\theta) - U(p, \mathcal{M}_i^\theta) \leq \|\tilde{p} - p\|_2 \|c\|_2. \quad (69)$$

Finally, Equations (68) and (69) together imply

$$|U(\tilde{p}, \mathcal{M}_i^\theta) - U(p, \mathcal{M}_i^\theta)| \leq \|\tilde{p} - p\|_2 \|c\|_2. \quad (70)$$

Equation (70) by definition says that U is Lipschitz-continuous in p with Lipschitz-constant $\|c\|_2$.

So, we finally obtain, that W is Lipschitz-continuous in p with Lipschitz-constant $(n+1)\|c\|_2$. $\square$

Lemma 4 (Convexity). *The map $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ from $\mathbb{R}_{\geq 0}^m$ to $\mathbb{R}_{\geq 0}$ is convex.¹³*

Proof. Since $W(p, (\mathcal{M}_i^\theta)_{i=1}^n) := R(p) + \sum_{i \in N} U(p, \mathcal{M}_i^\theta)$ is the sum of $1+n$ functions, we first quickly show that R is obviously convex in p and afterwards we will show that also U is convex in p .

¹²Note that Lipschitz-continuity implies (uniform) continuity. The Lipschitz-constant $(n+1)\|c\|_2$, given in the proof, only depends on the capacities c of $\mathcal{X}$ and on the number of bidders n . The proof also works for any other (possibly non-monotonic) value function $v_i : \mathcal{X} \rightarrow \mathbb{R}$ instead of $\mathcal{M}_i^\theta$.

¹³The proof also works for any other (possibly non-monotonic) value function $v_i : \mathcal{X} \rightarrow \mathbb{R}$ instead of $\mathcal{M}_i^\theta$.

From Equation (8) it follows that $R(p) = \sum_{j \in M} c_j p_j$ is linear in p and thus convex in p .

In the remainder of the proof we are going to show that $U(p, \mathcal{M}_i^\theta) := \max_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\}$ is convex in p . For every $i \in N$ and for every $x \in \mathcal{X}$, the map $p \mapsto \mathcal{M}_i^\theta(x) - \langle p, x \rangle$ is (affine-)linear¹⁴ in p and thus convex in p . As the maximum over convex functions is always convex (Danskin 1967; Bertsekas 1971, 1999), $U(p, \mathcal{M}_i^\theta) := \max_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\}$ is convex in p .

So finally we get that W is convex in p , since it is the sum of $(n+1)$ convex functions. $\square$

Lemma 5 (Sub-gradients). *Let $\hat{\mathcal{X}}_i^*(p) := \operatorname{argmax}_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\}$ denote the set of utility maximizing bundles w.r.t. $\mathcal{M}_i^\theta$ for the i -th bidder. Then the sub-gradients of the map $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ from $\mathbb{R}_{\geq 0}^m$ to $\mathbb{R}_{\geq 0}$ are given by*

$$\begin{aligned} & \nabla_p^{\text{sub}} W(p, (\mathcal{M}_i^\theta)_{i=1}^n) \\ &:= \operatorname{conv} \left\{ c - \sum_{i \in N} \hat{x}_i^*(p) : (\hat{x}_i^*(p))_{i \in N} \in \times_{i \in N} \hat{\mathcal{X}}_i^*(p) \right\}, \end{aligned} \quad (71)$$

where conv denotes the convex hull and $\times$ denotes the Cartesian product. In particular, for each $(\hat{x}_i^*(p))_{i \in N} \in \times_{i \in N} \hat{\mathcal{X}}_i^*(p)$,

$$c - \sum_{i \in N} \hat{x}_i^*(p) \in \nabla_p^{\text{sub}} W(p, (\mathcal{M}_i^\theta)_{i=1}^n) \quad (72)$$

is a subgradient of W with respect to p .¹⁵

Proof. Since $W(p, (\mathcal{M}_i^\theta)_{i=1}^n) := R(p) + \sum_{i \in N} U(p, \mathcal{M}_i^\theta)$ is the sum of $1 + n$ functions, we first quickly compute the (sub-)gradient of R with respect to p and afterwards we will compute the sub-gradients of U with respect to p .

From Equation (8) it follows that $R(p) = \sum_{j \in M} c_j p_j$ is linear in p and thus its sub-gradient is uniquely defined as its gradient $\nabla_p R(p) = c$, i.e., $\nabla_p^{\text{sub}} R(p) = \{c\}$.

In the remainder of the proof we are going to compute the set of sub-gradients of $U(p, \mathcal{M}_i^\theta) := \max_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\}$ with respect to p with the help of Danskin's theorem (originally proven by Danskin (1967) and later refined by Bertsekas (1971, 1999)). Let's define $\phi_i(p, x) := \mathcal{M}_i^\theta(x) - \langle p, x \rangle$. Then $U(p, \mathcal{M}_i^\theta) = \max_{x \in \mathcal{X}} \phi_i(p, x)$ and $\hat{\mathcal{X}}_i^*(p) = \operatorname{argmax}_{x \in \mathcal{X}} \phi_i(p, x)$. Next, we show that $\phi_i : \mathbb{R}_{\geq 0}^m \times \mathcal{X} \rightarrow \mathbb{R}$ fulfills all the assumptions of Danskin's theorem:

¹⁴In the following, we will call affine-linear functions "linear" too as commonly done in the literature.

¹⁵The proof also works for any other (possibly non-monotonic) value function $v_i : \mathcal{X} \rightarrow \mathbb{R}$ instead of $\mathcal{M}_i^\theta$.

- For every $x \in \mathcal{X}$ the map $p \mapsto \phi_i(p, x)$ is convex and differentiable, since it is (affine-)linear.
- For every $p \in \mathbb{R}_{\geq 0}^m$ the map $x \mapsto \phi_i(p, x)$ is continuous, since $\mathcal{X}$ is discrete and every function is continuous with respect to the discrete topology.
- Furthermore, the set $\mathcal{X}$ is compact (since it is finite).

Under these three assumptions Danskin's theorem tells us that

$$\nabla_p^{\text{sub}} U(p, \mathcal{M}_i^\theta) = \operatorname{conv} \left\{ \nabla_p \phi_i(p, x) : x \in \hat{\mathcal{X}}_i^*(p) \right\}. \quad (73)$$

Now we can simply compute $\nabla_p \phi_i(p, x)$ for any constant $x \in \mathcal{X}$, i.e.,

$$\nabla_p \phi_i(p, x) = \nabla_p (\mathcal{M}_i^\theta(x) - \langle p, x \rangle) \quad (74a)$$

$$= \nabla_p \mathcal{M}_i^\theta(x) - \nabla_p \langle p, x \rangle \quad (74b)$$

$$= 0 - x = -x. \quad (74c)$$

Plugging in Equation (74) into Equation (73) results in

$$\nabla_p^{\text{sub}} U(p, \mathcal{M}_i^\theta) = \operatorname{conv} \left\{ -x : x \in \hat{\mathcal{X}}_i^*(p) \right\}. \quad (75)$$

Note that set of sub-gradients of a sum of real-valued¹⁶ convex functions, is equal to the sum¹⁷ of the sets of sub-gradients of the individual functions (Rockafellar 1970, Theorem 23.8). Further note that the sum of convex hulls of sets is equal to the convex hull of their sum. We use these two insights and Equation (75) to finally compute the set of sub-gradients

$$\nabla_p^{\text{sub}} W(p, (\mathcal{M}_i^\theta)_{i=1}^n) = \quad (76a)$$

$$= \nabla^{\text{sub}} \left(R(p) + \sum_{i \in N} U(p, \mathcal{M}_i^\theta) \right) \quad (76b)$$

$$= \nabla^{\text{sub}} R(p) + \sum_{i \in N} \nabla^{\text{sub}} U(p, \mathcal{M}_i^\theta) \quad (76c)$$

$$= \{c\} + \sum_{i \in N} \operatorname{conv} \left\{ -x : x \in \hat{\mathcal{X}}_i^*(p) \right\} \quad (76d)$$

$$= \{c\} + \operatorname{conv} \sum_{i \in N} \left\{ -x : x \in \hat{\mathcal{X}}_i^*(p) \right\} \quad (76e)$$

$$= \{c\} - \operatorname{conv} \sum_{i \in N} \hat{\mathcal{X}}_i^*(p) \quad (76f)$$

$$= \{c\} - \operatorname{conv} \left\{ \sum_{i \in N} \hat{x}_i^*(p) : (\hat{x}_i^*(p))_{i \in N} \in \times_{i \in N} \hat{\mathcal{X}}_i^*(p) \right\} \quad (76g)$$

$$= \operatorname{conv} \left\{ c - \sum_{i \in N} \hat{x}_i^*(p) : (\hat{x}_i^*(p))_{i \in N} \in \times_{i \in N} \hat{\mathcal{X}}_i^*(p) \right\}, \quad (76h)$$

¹⁶Note that Rockafellar (1970, Theorem 23.8) is formulated for proper convex functions with overlapping effective domains that can potentially attain $+\infty$ as value. In our case of real-valued function (which implies that they cannot attain $+\infty$) all convex functions are proper and their effective domain is simply their domain.

¹⁷The (Minkowski) sum of two sets A and B is defined as the set $A + B = \{a + b : (a, b) \in A \times B\}$ of all possible sums of any element of the first set A and any element of the second set B .

which proves the main statement (71) of Lemma 5. This directly implies that for each $(\hat{x}_i^*(p))_{i \in N} \in \times_{i \in N} \hat{\mathcal{X}}_i^*(p)$,

$$c - \sum_{i \in N} \hat{x}_i^*(p) \in \nabla_p^{\text{sub}} W(p, (\mathcal{M}_i^\theta)_{i=1}^n) \quad (77)$$

is an element of the convex hull in (76h) and thus a sub-gradient of W with respect to p . $\square$

Next, Lemma 6 shows that the map $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ from $\mathbb{R}_{\geq 0}^m$ to $\mathbb{R}_{\geq 0}$ is Lebesgue-almost everywhere differentiable with (proper) gradient¹⁸

$$\nabla_p W(p, (\mathcal{M}_i^\theta)_{i=1}^n) \stackrel{\text{a.e.}}{=} c - \sum_{i \in N} \hat{x}_i^*(p), \quad (78)$$

where $\hat{x}_i^*(p) \in \hat{\mathcal{X}}_i^*(p) := \operatorname{argmax}_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\}$ denotes the utility maximizing bundle w.r.t. $\mathcal{M}_i^\theta$ for the i -th bidder, which is Lebesgue-almost everywhere unique.

Lemma 6 (A.e. Differentiable). *Let $(\mathcal{M}_i^\theta)_{i=1}^n$ be a tuple of m MVNNs. Then there exists a dense¹⁹ subset $P \subseteq \mathbb{R}_{\geq 0}^m$ such that $\mathbb{R}_{\geq 0}^m \setminus P$ is a Lebesgue null set²⁰, and that $\forall p \in P$:*

1. A unique optimizer $\hat{x}_i^*(p) \in \hat{\mathcal{X}}_i^*(p)$ exists, i.e.,

$$\{x_i^*(p)\} = \hat{\mathcal{X}}_i^*(p) := \operatorname{argmax}_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\} \quad (79)$$

and

2. the map $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ is differentiable with gradient

$$\nabla_p W(p, (\mathcal{M}_i^\theta)_{i=1}^n) = c - \sum_{i \in N} \hat{x}_i^*(p), \quad (80)$$

which is also the unique sub-gradient, i.e. $\{c - \sum_{i \in N} \hat{x}_i^*(p)\} = \nabla_p^{\text{sub}} W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$.

3. The unique optimizer $\hat{x}_i^*(p)$ is constant in a local neighborhood of p , and thus $\partial_p \hat{x}_i^*(p) = 0$.²¹

Proof. We start this proof by defining P and showing a.e. differentiability. Since $W(p, (\mathcal{M}_i^\theta)_{i=1}^n) := R(p) + \sum_{i \in N} U(p, \mathcal{M}_i^\theta)$ is the sum of $1 + n$ functions, we first quickly show that R is obviously differentiable and afterwards we show that each $U(p, \mathcal{M}_i^\theta)$ is differentiable on a set P_i . Afterwards we will define P as their intersection.

From Equation (8) it follows that $R(p) = \sum_{j \in M} c_j p_j$ is linear in p and thus differentiable.

¹⁸Note that the gradient in Equations (78) and (80) actually denotes a classical gradient and not just a sub-gradient.

¹⁹The set P being *dense* means that its topological closure $\bar{P}$ covers the whole space, i.e., $\bar{P} = \mathbb{R}_{\geq 0}^m$.

²⁰The set $\mathbb{R}_{\geq 0}^m \setminus P$ being a (*Lebesgue*) *null set* means that its m -dimensional Lebesgue measure (i.e., the m -dimensional volume) $\lambda^m(\mathbb{R}_{\geq 0}^m \setminus P) = 0$. Within this work “a.e.” always corresponds to “Lebesgue almost everywhere”.

²¹Note that one could even prove that the gradient is continuous on P , while the gradient can have jumps at the null set $\mathbb{R}_{\geq 0}^m \setminus P$. The proof also works for any other (possibly non-monotonic) value function $v_i : \mathcal{X} \rightarrow \mathbb{R}$ instead of $\mathcal{M}_i^\theta$.

Next, we show that $U(p, \mathcal{M}_i^\theta) := \max_{x \in \mathcal{X}} \{\mathcal{M}_i^\theta(x) - \langle p, x \rangle\}$ is differentiable a.e. and define P_i . We know already from the proof of Lemma 4 that $U(p, \mathcal{M}_i^\theta)$ is convex and Rockafellar (1970, Theorem 25.5 on p. 246) tells us that for any²² real-valued convex function there exists a dense set P_i on which the functions is differentiable with the complement of P_i being a null set. Thus, for each $p \mapsto U(p, \mathcal{M}_i^\theta)$, $i \in N$, we obtain such a P_i . If $1 + n \in \mathbb{N}$ functions are differentiable at a point p , then their sum is too. Thus, $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ is differentiable on $P := \bigcap_{i=1}^n P_i$. Since each P_i is dense in $\mathbb{R}_{\geq 0}^m$, P is also dense in $\mathbb{R}_{\geq 0}^m$. Moreover, it holds that $\lambda^m(\mathbb{R}_{\geq 0}^m \setminus P) = \lambda^m(\bigcup_{i=1}^n (\mathbb{R}_{\geq 0}^m \setminus P_i)) \leq \sum_{i=1}^n \lambda^m(\mathbb{R}_{\geq 0}^m \setminus P_i) = \sum_{i=1}^n 0 = 0$, i.e., the m -dimensional Lebesgue measure of $\mathbb{R}_{\geq 0}^m \setminus P$ vanishes. Putting everything together, we get that $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ is a.e. differentiable (concretely differentiable for all $p \in P$).

Next, we prove the uniqueness of $\hat{x}_i^*(p) \in \hat{\mathcal{X}}_i^*(p)$ for every $p \in P_i \supseteq P$. We know that $p \mapsto U(p, \mathcal{M}_i^\theta)$ is differentiable at $p \in P_i$, and thus we know that the sub-gradient is unique, i.e., $|\nabla_p^{\text{sub}} U(p, \mathcal{M}_i^\theta)| = |\{\nabla_p U(p, \mathcal{M}_i^\theta)\}| = 1$. We have already computed this set of sub-gradients $\nabla_p^{\text{sub}} U(p, \mathcal{M}_i^\theta) = \operatorname{conv} \{-x : x \in \hat{\mathcal{X}}_i^*(p)\} = \operatorname{conv} -\hat{\mathcal{X}}^*(p) \supseteq -\hat{\mathcal{X}}^*(p)$ in Equation (75) in the proof of Lemma 5. This set can only be a singleton, if $\hat{\mathcal{X}}^*(p)$ is a singleton. Finally, $|\hat{\mathcal{X}}^*(p)| = 1$ immediately implies Item 1, i.e., the uniqueness of $\hat{x}_i^*(p) \in \hat{\mathcal{X}}_i^*(p)$.

Using that $p \mapsto U(p, \mathcal{M}_i^\theta)$ is differentiable for every $p \in P$ together with Item 1 and Lemma 5, we finally obtain Item 2 for every $p \in P$.

For Item 3 it is crucial that $\mathcal{X}$ is finite. For every $p \in P$, we know from Item 1 that there is a unique maximizer $\hat{x}_i^*(p) \in \hat{\mathcal{X}}_i^*(p) = \operatorname{argmax}_{x \in \mathcal{X}} \phi_i(p, x)$ with $\phi_i(p, x) := \mathcal{M}_i^\theta(x) - \langle p, x \rangle$.

Since $\mathcal{X}$ is finite, we can define

$$\epsilon_i(p) := \phi_i(p, \hat{x}_i^*(p)) - \max_{x \in \mathcal{X} \setminus \{\hat{x}_i^*(p)\}} \phi_i(p, x), \quad (81)$$

which has to be strictly larger than 0 for every $p \in P$ because of the uniqueness of the maximizer. Equation (81) implies that $\hat{x}_i^*(p)$ outperforms every other $x \in \mathcal{X}$ by at least a margin of $\epsilon_i(p)$. Since ϕ_i is continuous, $\hat{x}_i^*(p)$ cannot be “overtaken” by any other $x \in \mathcal{X}$ within a small neighbourhood of p as we will calculate explicitly in the following using the Lipschitz constant $\|c\|_2$ of ϕ_i that we derived in the proof of

²²There are some very mild technical assumptions in Rockafellar (1970, Theorem 25.5 on p. 246) that do not matter in our case. Rockafellar (1970) assumes that the function is defined on $\mathbb{R}^m$. In our case we could easily extend the domain of our function from $\mathbb{R}_{\geq 0}^m$ to $\mathbb{R}^m$. As Rockafellar (1970, Theorem 25.5 on p. 246) is formulated for proper convex functions that can also attain $+\infty$ one could extend any convex function from a convex domain (such as $\mathbb{R}_{\geq 0}^m$) to be defined to be $+\infty$ outside of that convex domain. Note that if P_i is dense in $\mathbb{R}_{\geq 0}^m$ it is also dense in $\mathbb{R}_{\geq 0}^m$ and that $\lambda^m(\mathbb{R}_{\geq 0}^m \setminus \mathbb{R}_{\geq 0}^m) = 0$.

Lemma 3: For any $x \in \mathcal{X} \setminus \{\hat{x}_i^*(p)\}$, $p \in P$, $\tilde{p} \in \mathbb{R}_{\geq 0}^m$, we have

$$\phi_i(\tilde{p}, x) \leq \phi_i(p, x) + \|c\|_2 \|\tilde{p} - p\|_2 \quad (82a)$$

$$\leq \max_{x \in \mathcal{X} \setminus \{\hat{x}_i^*(p)\}} \phi_i(p, x) + \|c\|_2 \|\tilde{p} - p\|_2 \quad (82b)$$

$$\leq \phi_i(p, \hat{x}_i^*(p)) - \epsilon_i(p) + \|c\|_2 \|\tilde{p} - p\|_2 \quad (82c)$$

$$\leq \phi_i(\tilde{p}, \hat{x}_i^*(p)) - \epsilon_i(p) + 2 \|c\|_2 \|\tilde{p} - p\|_2. \quad (82d)$$

From this inequality we obtain, that within an open ball with radius $\min_{i \in N} \frac{\epsilon_i(p)}{2\|c\|_2} > 0$ around $p \in P$, the optimizers $\hat{x}_i^*(p)$ stay constant for every $i \in N$. Therefore, the differential $\partial_p \hat{x}_i^*(p)$ is zero for all $p \in P$, which concludes the proof of Item 3. $\square$

Putting everything together, we can finally prove Theorem 3.

Proof of Theorem 3. Combining Lemmas 3 to 6 proves Theorem 3 (and even slightly stronger statements, e.g., Lemma 5 fully specifies the set of all sub-gradients). $\square$

Theorem 3 and most of the statements from Lemmas 3 to 6 can be proven under even less assumptions as we will discuss in the following Remarks C.1 and C.2 which generalize the theory to further settings.

Remark C.1 (General Value Functions). *Theorem 3 and Lemmas 3 to 6 are also true for the map $p \mapsto W(w, (g_i)_{i=1}^n)$, with any (possibly non-monotonic) value function $g_i : \mathcal{X} \rightarrow \mathbb{R}$ instead of $g_i = \mathcal{M}_i^\theta$, since we never used any specific properties of MVNNs $\mathcal{M}_i^\theta$ in the proofs of these statements. In particular, this includes bidders' true value functions $g_i = v_i$, $v_i : \mathcal{X} \rightarrow \mathbb{R}$.*

Remark C.2 (Continuous Input Space $\tilde{\mathcal{X}}$). *Theorem 3 and all of the statements from Lemmas 3 to 6 except Item 3 from Lemma 6 are also true if one replaces the finite set $\mathcal{X}$ by any (possibly non-discrete) compact set $\tilde{\mathcal{X}}$ (if one extends the definition of $\mathcal{M}_i^\theta$ in the natural way from $\mathcal{X}$ to $\tilde{\mathcal{X}} \subseteq \mathbb{R}^m$). When combining Remarks C.1 and C.2 one has to assume that the $g_i : \tilde{\mathcal{X}} \rightarrow \mathbb{R}$ are continuous for our proof.²³ Item 3 from Lemma 6 can be violated for compact $\tilde{\mathcal{X}} \subseteq \mathbb{R}^m$ if $|\tilde{\mathcal{X}}| = \infty$.²⁴*

Remark C.3 (Piece-wise Linear). *In our case of finite $\mathcal{X}$, one can intuitively see with similar arguments as in the proof of Lemma 4, that $p \mapsto W(p, (\mathcal{M}_i^\theta)_{i=1}^n)$ is piece-wise linear as a sum over a linear function and n functions which are the maxima over finitely many linear functions.*

D Experiment Details

In this section, we present all details of our experiments from Section 6.

²³Note that if $\tilde{\mathcal{X}}$ is finite, every function $g_i : \tilde{\mathcal{X}} \rightarrow \mathbb{R}$ is continuous by definition.

²⁴Note that, while Item 3 from Lemma 6 was used for the intuitive sketch of the proof of Theorem 3 in the main paper, Item 3 from Lemma 6 is not necessary at all for the mathematical rigorous proof of Theorem 3 given in Appendix C.2.

D.1 SATS Domains

In this section, we provide a more detailed overview of the four SATS domains, which we use to experimentally evaluate ML-CCA:

- **Global Synergy Value Model (GSVM)** (Goeree and Holt 2010) has 18 items with capacities $c_j = 1$ for all $j \in \{1, \dots, 18\}$, 6 *regional* and 1 *national bidder*. In GSVM the value of a package increases by a certain percentage with every additional item of interest. Thus, the value of a bundle only depends on the total number of items contained in a bundle which makes it one of the simplest models in SATS. In fact, bidders' valuations exhibit at most two-way (i.e., pairwise) interactions between items.
- **Local Synergy Value Model (LSVM)** (Scheffell, Ziegler, and Bichler 2012) has 18 items with capacities $c_j = 1$ for all $j \in \{1, \dots, 18\}$, 5 *regional* and 1 *national bidder*. Complementarities arise from spatial proximity of items.
- **Single-Region Value Model (SRVM)** (Weiss, Lubin, and Seuken 2017) has 3 items with capacities $c_1 = 6, c_2 = 14, c_3 = 9$ and 7 bidders (categorized as *local*, *high frequency*, *regional*, or *national*) and models UK 4G spectrum auctions.
- **Multi-Region Value Model (MRVM)** (Weiss, Lubin, and Seuken 2017) has 42 items with capacities $c_j \in \{2, 3\}$ for all $j \in \{1, \dots, 42\}$ and 10 bidders (*local*, *regional*, or *national*) and models large Canadian 4G spectrum auctions.

In the efficiency experiments in this paper, we instantiated for each SATS domain the 100 synthetic CA instances with the seeds $\{101, \dots, 200\}$. We used SATS version 0.8.1.

D.2 Compute Infrastructure

All experiments were conducted on a compute cluster running Debian GNU/Linux 10 with Intel Xeon E5-2650 v4 2.20GHz processors with 24 cores and 128GB RAM and Intel E5 v2 2.80GHz processors with 20 cores and 128GB RAM and Python 3.8.10.

D.3 Hyperparameter Optimization

In this section, we provide details on our exact HPO methodology and the ranges that we used.

We separately optimized the HPs of the mMVNNs for each bidder type of each domain, using a different set of SATS seeds than for all other experiments in the paper. Specifically, for each bidder type, we first trained an mMVNN using as initial data points the demand responses of an agent of that type during 50 consecutive CCA clock rounds, and then measured the generalization performance of the resulting network on a validation set that was created by drawing 500 price vectors where the price of each item was drawn uniformly at random from the range of zero to three times the average maximum value of an agent of that type for a single item (which was determined using separate seeds, see validation set 2 in Figure 1). The number of seeds used to evaluate each model was equal for all models and set to 10. Finally, for each bidder type we selected the set of HPs that performed the best on this validation set with respect to the coefficient of determination (R^2). The full range of HPs tested for all agent types and all domains is shown

Hyperparameter	HPO-Range
Non-linear Hidden Layers	[1,2,3]
Neurons per Hidden Layer	[8, 10, 20, 30]
Learning Rate	(1e-4, 1e-2)
Epochs ²⁵	[30, 50, 70, 100]
L2-Regularization	(1e-8, 1e-2)
Linear Skip Connections ²⁶	[True, False]

Table 2: HPO ranges for all domains.

in Table 2, while the winning configurations are shown in Table 3.

Additionally, we determined the set of HPs with the best generalization performance on validation set 2 using as evaluation metric a shift-invariant variation of R^2 , defined as:

$$R_c^2 = 1 - \frac{\sum_r ((v_i(x^r) - \bar{v}_i) - (\mathcal{M}_i(x^r) - \bar{\mathcal{M}}_i))^2}{\sum_r (v_i(x^r) - \bar{v}_i)^2}, \quad (83)$$

where $v_i(x^r)$ is the true value of the bidder for the r -th bundle, $\mathcal{M}_i(x^r)$ is the neural network’s predicted value for that bundle, and $\bar{v}_i$ and $\bar{\mathcal{M}}_i$ are their empirical means, respectively. The reason that we opted for this shift-invariant version of R^2 is that, as explained in Section 3.1, learning the true value functions of the agents up to a constant shift suffices for our query generation procedure as described in Section 4. Surprisingly, in all domains our mechanism performed slightly worse with those HPs, with the maximum efficiency delta between the two configurations being 1.2% in LSVM. However, in all domains results were qualitatively identical. The winning configurations for both metrics are shown in Tables 3 and 4. In all domains we chose the configurations from Table 3 for our efficiency experiments.

D.4 Details on mMVNN Training

Remark D.1 (Other ML-models). *Note that our training method TRAINONDQS (Algorithm 1) also works for any other ML method that can be trained via GD and for which the inner optimization problem $\hat{x}_i^*(p^r) \in \arg\max_{x \in \mathcal{X}} \{\mathcal{M}_i^{\theta_t}(x) - \langle p^r, x \rangle\}$ (see Line 4 of Algorithm 1) can be solved efficiently. This is the case for (m)MVNNs, where Line 4 can be solved as a MILP analogously to (Weissteiner et al. 2022a). Another example would be classical ReLU-neural networks (NNs)²⁷, where such a MILP formulation exists too (Weissteiner and Seuken 2020), which are suitable for domains without the free disposal property.*

Remark D.2 (Initialization). *We use the initialization scheme introduced by Weissteiner et al. (2023), which offers*

²⁵For GSVM and LSVM, the number of epochs was fixed to 30

²⁶For the definition of (m)MVNNs with a linear skip connection, please see Weissteiner et al. (2023, Definition F.1)

²⁷Note that in principal for every NN with piece-wise linear activation function (e.g., ReLU, bReLU or Leaky ReLU) a MILP-formulation is possible. However for other activation functions such as the sigmoid activation function an exact MILP-formulation is not possible.

advantages over the original initialization scheme used by Weissteiner et al. (2022a) as explained in Weissteiner et al. (2023, Section 3.2 and Appendix E).

In the conducted experiments, Python 3.8.10 and PyTorch 2.0.0, were employed as the primary programming language and framework for implementing the mMVNNs. The Adam optimizer was chosen as the optimization algorithm for the training process. To further enhance the training procedure, the cosine annealing scheduler was utilized, dynamically adjusting the learning rate over epochs to facilitate convergence and prevent premature stagnation.

D.5 Details MILP Parameters

There are three distinct points in which MILPs are solved in our ML-powered combinatorial clock auction (ML-CCA):

1. in the training of mMVNNs according to Line 4 in Algorithm 1,
2. for W minimization according to Theorem 3 in order to predict the demand of each agent at a given price vector (see Line 9 in Algorithm 3), and
3. finally to solve the *winner determination problem (WDP)* and determine the resulting allocation based on the elicited bids (see Line 18 in Algorithm 2).

The first two MILPs are of the same type: given as input an mMVNN $\mathcal{M}_i^{\theta}$ that approximates a bidder’s value function and linear item prices $p \in \mathbb{R}_{\geq 0}^m$, find the utility maximizing bundle for that bidder, i.e., solve $\hat{x}_i^*(p) \in \arg\max_{x \in \mathcal{X}} \{\mathcal{M}_i^{\theta}(x) - \langle p, x \rangle\}$. The third MILP is of a different type: given as inputs a set of bundle-value tuples from each agent, find a feasible allocation that maximizes reported social welfare. This WDP is described in more detail in Section 2 and Equation (2). In each clock round, only two WDP MILPs need to be solved, one involving just the clock bids of the agents and one also including the bids that would result from the clock bids raised heuristic. For each agent, on average two thousand MILPs of type 1 need to be solved per clock round. It should be noted that they are very fast, as a single MILP of this type can be solved in under 200 milliseconds in our server architecture as described in Appendix D.2. The MILPs of the first type used a formulation based on the MILP-formulation for MVNNs in (Weissteiner et al. 2023, Section 3.2 and Appednix F) which is an improved version of the MILP-formulation for MVNNs in (Weissteiner et al. 2022a, Theorem 2 in Section 3.1 and Appendix C.5). The MILPs of the first type were solved using the Gurobi 10 solver, and the WDP MILPs were solved using CPLEX 20.01.0. For all MILPs, we set the feasibility tolerance to $1e-9$, the integrality tolerance to $1e-8$ and the relative MIP optimality gap to $1e-06$. All other parameters were set to their respective default values.

D.6 Details on NEXTPRICE Procedure

In this section, we describe the details of the NEXTPRICE procedure from Line 9 in Algorithm 2. Given trained mMVNNs, NEXTPRICE generates new demand queries by minimizing W under constraint (9b) via GD which is based on Theorems 2 and 3.

DOMAIN	BIDDER TYPE	# HIDDEN LAYERS	# HIDDEN UNITS	LIN. SKIP	LEARNING RATE	L2 REGULARIZATION	EPOCHS
GSVM	REGIONAL	2	20	FALSE	0.005	0.00001	30
	NATIONAL	3	30	TRUE	0.001	0.000001	30
LSVM	REGIONAL	1	30	TRUE	0.01	0.000001	30
	NATIONAL	3	20	FALSE	0.005	0.0001	30
SRVM	LOCAL	2	20	TRUE	0.01	0.0001	30
	REGIONAL	1	20	TRUE	0.01	0.0001	50
	NATIONAL	1	30	FALSE	0.005	0.00001	70
	HIGH FREQUENCY	2	20	FALSE	0.01	0.00001	30
MRVM	LOCAL	3	20	TRUE	0.005	0.000001	100
	REGIONAL	2	20	TRUE	0.001	0.001	100
	NATIONAL	3	20	TRUE	0.001	0.0001	50

Table 3: Winning HPO configurations for R^2

DOMAIN	BIDDER TYPE	# HIDDEN LAYERS	# HIDDEN UNITS	LIN. SKIP	LEARNING RATE	L2 REGULARIZATION	EPOCHS
GSVM	NATIONAL	3	10	TRUE	0.001	0.001	30
	REGIONAL	2	10	TRUE	0.01	0.001	30
LSVM	NATIONAL	1	10	TRUE	0.005	0.01	30
	REGIONAL	3	20	TRUE	0.005	0.0001	30
SRVM	LOCAL	2	30	TRUE	0.01	0.0001	30
	REGIONAL	2	20	TRUE	0.01	0.000001	50
	NATIONAL	1	20	TRUE	0.01	0.0001	50
	HIGH FREQUENCY	2	30	TRUE	0.01	0.00001	70
MRVM	NATIONAL	1	20	TRUE	0.001	0.000001	30
	REGIONAL	3	20	TRUE	0.001	0.000001	50
	LOCAL	2	20	FALSE	0.001	0.000001	30

Table 4: Winning HPO configurations for R_c^2

Detailed motivation of constraint (9b) In the case that linear clearing prices (LCPs) exist, both minimizing W with or without constraint (9b) would lead to clearing prices and thus to efficient allocations, if the mMVNNs approximate the value functions v well enough, as shown in Theorem 2 and Lemma 2, respectively. In this case, constraint (9b) is neither beneficial nor harmful, because both versions are well motivated by theory (see Lemma 2 and Theorem 2): In this case the set of solutions to both problems (9) and (37) are both exactly equal to the set of all possible LCPs and thus result in efficient allocations with no over-demand and no under-demand (see the proof of Lemma 2 and Theorem 2 in Appendix C.1). Thus in the case that LCPs do exist, both problems (9) and (37) are exactly equivalent and well supported by theory. Indeed our experiments resulted in similarly good results for both versions of the method in domains where LCPs often exist (see GSVM, LSVM, SRVM in Tables 10 and 11).²⁸

²⁸In Tables 10 and 11 one can see a slight tendency that also for GSVM, LSVM, SRVM adding constraint (9b) is rather beneficial on average. The reason for this might be, that we do not always find LCPs even in these domains (note that LCPs do not always exist in

However, when no LCPs exist, for every price vector p we have over-demand or under-demand for some goods and we need to make a choice on how to select a price vector p for the next demand query based on the over- and under-demand.

Minimizing the classical W introduced in Lemma 2 via a classical GD-update rule using the gradient derived in Theorem 3 punishes over- and under-demand symmetrically. This can be directly seen from the classical GD-update rule

$$p_j^{\text{new}} \stackrel{\text{a.e.}}{=} p_j - \gamma(c_j - \sum_{i \in N} (\hat{x}_i^*(p))_j), \forall j \in M,$$

since the magnitude of change in the price vector would be the same for the same amount of over- or under-demand.

The following example also illustrates this symmetry.

Example D.1. Suppose there is a single item, and two bidders with a value of 5 and $5 - \epsilon$ for that item. Any price $p \in [5 - \epsilon, 5]$ is a clearing price, where the indirect utility of the bidder with the higher value is $5 - p$ and of the bidder with the lower value is 0, while the seller's indirect revenue

these domains).

is p for a W value of 5. For a price that is $x \in \mathbb{R}_{>0}$ higher than the largest clearing price, i.e., $5 + x$, no agent buys the item and they have an indirect utility of 0, while the seller's indirect revenue is $5 + x$, for a W of value $5 + x$. For a price that is x lower than the smallest clearing price, i.e., $5 - \epsilon - x$, both agents want to buy the item and they have indirect utilities of $\epsilon + x$ and x , while the seller's indirect revenue is $5 - \epsilon - x$, for a total W value of $5 + x$.

However, even though the W objective of Lemma 2 punishes over- and under-demand equally, our preference between them is highly asymmetric; we strongly prefer under-demand over over-demand in practice, since the demand responses of the agents at a price vector with no over-demand constitute a feasible allocation, while the demand responses of the agents at a price vector with over-demand do not. This is important because in case that the market does not clear within the clock round limit, our ML-CCA, just like the CCA, will have to combine the clock bids of the agents to produce a *feasible* allocation with the highest inferred social welfare according to Equation (2). If the demand responses elicited from the agents constituted feasible solutions, it makes it more likely that they can be effectively combined together in the WDP of Equation (2) to produce highly efficient allocations. This is why in domains where no LCPs exist, adding constraint (9b) leads to significantly increased efficiency (see MRVM in Table 11). For more intuition on constraint (9b) see the following example.

Example D.2. Suppose there are $m = 2$ items with capacities $c_1 = c_2 = 10$ and $n = 2$ bidders with value functions

$$v_1(x) = \max \{10\mathbb{1}_{\{x \geq (7,3)\}}, 10\mathbb{1}_{\{x \geq (3,7)\}}, 9\mathbb{1}_{\{x \geq (4,4)\}}\}, \\ v_2(x) = \max \{10\mathbb{1}_{\{x \geq (8,2)\}}, 10\mathbb{1}_{\{x \geq (2,8)\}}, 9\mathbb{1}_{\{x \geq (4,4)\}}\}.$$

In this setting, no LCP exists. This can be seen as follows. First note that we can obviously exclude every price vector $p \in \mathbb{R}_{\geq 0}^2$ with $p_1 = 0$ or $p_2 = 0$ from being a LCP. Furthermore, for any price vector $p \in \mathbb{R}_{>0}^2$ bidder 1's utility maximizing bundle $x_1^*(p) \in \mathcal{X}_1^*(p) = \{(7, 3), (3, 7), (4, 4)\}$ and bidder 2's utility maximizing bundle $x_2^*(p) \in \mathcal{X}_2^*(p) = \{(8, 2), (2, 8), (4, 4)\}$. However, from this we see that $x_1^*(p)$ and $x_2^*(p)$ cannot be combined without over- or under-demand, thus violating Item 2 in Definition 3.

The clearing potential objective W is minimized for every price vector $p = (p_1, p_2)$ that satisfies $p_1 = p_2 \in [0, 0.5]$.²⁹ For $p_1 = p_2 \in (0, 0.5)$, we have $\mathcal{X}_1^*(p) = \{(7, 3), (3, 7)\}$ and $\mathcal{X}_2^*(p) = \{(8, 2), (2, 8)\}$ and thus there is always positive over-demand for at least one of the items, which violates constraint (9b) (this is also the case for $p = (0, 0)$). For $p = (0.5, 0.5)$, we have $\mathcal{X}_1^*(p) = \{(7, 3), (3, 7), (4, 4)\}$ and $\mathcal{X}_2^*(p) = \{(8, 2), (2, 8), (4, 4)\}$. Thus, $p = (0.5, 0.5)$ fulfills constraint (9b), since $((4, 4), (4, 4)) \in \mathcal{X}_1^*(p) \times \mathcal{X}_2^*(p)$

²⁹The objective function W can be formulated as $W(p) =$

$$\begin{aligned} & \max \{0, 10 - \langle (7, 3), p \rangle, 10 - \langle (3, 7), p \rangle, 9 - \langle (4, 4), p \rangle\} \\ & + \max \{0, 10 - \langle (8, 2), p \rangle, 10 - \langle (2, 8), p \rangle, 9 - \langle (4, 4), p \rangle\} \\ & + (10x + 10y) \quad \forall p \in \mathbb{R}_{\geq 0}^2. \end{aligned}$$

The set of minimizers can be seen by plotting W .

is feasible (see Footnote 3). Thus, $p = (0.5, 0.5)$ is the unique solution of the constrained problem (9). In this case, $((4, 4), (4, 4))$ would be the efficient allocation with a SCW of 18. If we had only asked demand queries for prices $p_1 = p_2 \in [0, 0.5)$, i.e., prices that solve the unconstrained minimization problem (37), but not the constrained minimization problem (9), the WDP would end up with a SCW of only 10 by allocating some bundle of value 10 to one of the bidders and nothing to the other bidder, since the WDP is constrained by feasibility due to the limited capacity $c = (10, 10)$.

Details on NEXTPRICE (Algorithm 3) In Algorithm 3, we present the details of our NEXTPRICE procedure, a modification of the classical GD in Theorem 3 that systematically favours under-demand over over-demand (see Section 4). Compared to classical gradient descent on W (based on Theorem 3), there are three noteworthy modifications in the NEXTPRICE procedure.

1. First, as outlined at the end of Section 4, to incentivize GD on $W(p, (\mathcal{M}_i^0)_{i=1}^n)$ towards price vectors with no positive over-demand, the GD steps punish over- and under-demand *asymmetrically*. Specifically, at each iteration step, in case of predicted over-demand for some good, the gradient step for that good is $(1 + \mu)$ -times larger than what it would have been in case of under-demand (Line 26). Finally, μ is not a constant, but it adaptively increases as long as Algorithm 3 has not found a price vector with no predicted over-demand (Line 29).
2. Second, as also outlined at the end of Section 4, once the gradient steps have terminated, Algorithm 3 returns the price vector p that led to the lowest value of W among all price vectors examined that led to no positive predicted over-demand (i.e., satisfying constraint (9b)) (Lines 16 to 18).³⁰
3. Finally, the learning rate for each good in Algorithm 3 is scaled to be proportional to that good's current price p_j (Line 23). In theory, we do not have to do this, as Theorem 3 guarantees that even a uniform learning rate for all goods cannot get stuck in local minima of W . Using a uniform learning rate for all goods has two disadvantages. First, we would have to tune that learning rate parameter separately for each domain, since the goods' values are in different scales in each domain. Additionally, in the SRVM and MRVM domains the prices of different goods can vary by orders of magnitude. If we were to select a uniform learning rate for all goods, we would have to select one that would be suitable for the lowest-valued items (otherwise the GD steps would overshoot a lot for the prices of lower-valued goods, i.e. GD, would jump back and forth between large amounts of over- and under-demand for lower-valued goods), which would increase significantly the number of steps required until the learning rate becomes sufficiently small so that we do

³⁰If every gradient step resulted in positive predicted over-demand, we would pick just the one with minimal W ignoring the constraint (9b) for this demand query (Lines 30 to 31), but this case never occurred in our experiments (see Figure 3).

Algorithm 3: NEXTPRICE

Input : Trained MVNNs $(\mathcal{M}_i^\theta)_{i=1}^n$, last F_{init} round prices $p_j^{Q_{\text{init}}}$, Epochs $T \in \mathbb{N}$, Learning Rate Base $\lambda > 0$, Learning Rate Decay $\eta \in [0, 1]$, Feasibility multiplier μ , Feasibility multiplier increment ν

```

1 for  $j = 1$  to  $m$  do
2    $p_j^0 \leftarrow \sim U[0.75 \cdot p_j^{Q_{\text{init}}}, 1.25 \cdot p_j^{Q_{\text{init}}}]$ 
3    $W_{\text{best}} \leftarrow \infty$ 
4    $W_{\text{best}}^f \leftarrow \infty$ 
5   feasible  $\leftarrow$  False
6   for  $t = 0$  to  $T - 1$  do
7      $W \leftarrow \langle p^t, c \rangle$   $\triangleright$  Seller's indirect revenue
8     for  $i = 1$  to  $n$  do
9       Solve  $\hat{x}_i^*(p^t) \in \operatorname{argmax}_{x \in \mathcal{X}} \mathcal{M}_i^\theta(x) - \langle p^t, x \rangle$ 
10       $U_i \leftarrow \mathcal{M}_i^\theta(\hat{x}_i^*(p^t)) - \langle p^t, \hat{x}_i^*(p^t) \rangle$   $\triangleright$  Bidder  $i$ 's indirect Utility
11       $W \leftarrow W + U_i$ 
12       $d \leftarrow \sum_{i \in N} \hat{x}_i^*(p^t)$   $\triangleright$  Total Predicted Demand
13      if  $W < W_{\text{best}}$  then  $\triangleright$  Found better prices wrt.  $W$ 
14         $W_{\text{best}} \leftarrow W$ 
15         $p_{\text{best}} \leftarrow p^t$ 
16      if  $W < W_{\text{best}}^f$  and  $d \leq c$  then  $\triangleright$  Found better feasible prices
17         $W_{\text{best}}^f \leftarrow W$ 
18         $p_{\text{best}}^f \leftarrow p^t$ 
19        feasible  $\leftarrow$  True
20      if  $d = c$  then  $\triangleright$  Predicted Market Clearing Prices
21        break
22      for  $j = 1$  to  $m$  do
23         $\gamma_j \leftarrow \lambda \cdot p_j^t$   $\triangleright$  Scale l.r. for each good
24         $p_j^{t+1} \leftarrow p_j^t - \gamma_j(c_j - d_j)$   $\triangleright$  Theorem 3, eq. (11a)
25        if  $d_j > c_j$  then  $\triangleright$  Over-demand for good  $j$ 
26           $p_j^{t+1} \leftarrow p_j^t - \mu \gamma_j(c_j - d_j)$   $\triangleright$  Equation (11)
27           $\lambda \leftarrow \lambda \cdot (1 - \eta)$   $\triangleright$  Learning rate decay
28        if not feasible then  $\triangleright$  No feasible allocation yet
29           $\mu \leftarrow \mu + \nu$ 
30      if not feasible then  $\triangleright$  No feasible allocation found
31         $p_{\text{best}}^f \leftarrow p_{\text{best}}$   $\triangleright$  Footnote 30
32      if  $\mu = \nu = 0$  then  $\triangleright$  Do not enforce feasibility (see Remark D.3)
33      return Prices  $p_{\text{best}}$  minimizing  $W(\cdot, (\mathcal{M}_i^\theta)_{i=1}^n)$ 
34 return Feasible prices  $p_{\text{best}}^f$  minimizing  $W(\cdot, (\mathcal{M}_i^\theta)_{i=1}^n)$ 

```

not have such extreme jumps. Scaling the learning rate for each good proportionally to its current price alleviates both of these potential issues.

Remark D.3. Note that by setting $\mu = \nu = 0$, NEXTPRICE (Algorithm 3) performs symmetrical GD on W without constraint (9b) as suggested by Theorem 3 (see Appendix D.9 for an empirical evaluation of minimizing W with $\mu = \nu = 0$, i.e., without constraint (9b)).

Remark D.4. For each GD-step we solve the inner optimization problem of Equation (7), i.e., $\max_{x \in \mathcal{X}} (\mathcal{M}_i^\theta(x) - \langle p, x \rangle)$, in Line 9 for each bidder i ,

using the MILP encoding of MVNNs from (Weissteiner et al. 2022a).³¹

In our experiments, we use 300 epochs for all domains, with a good-specific learning rate of 1% of the price p_j^t of that good and a learning rate decay of 0.5%, i.e., we set $T = 300$, $\lambda = 0.01$ and $\eta = 0.005$, while we set $\mu = 2$ and $\nu = 1.01$.

Intuitively, this way Algorithm 3 punishes over-demand at least three times as much as under-demand, which means that it can very quickly converge to a price region with no over-demand, and then it starts minimizing under-demand in the same way as suggested by Theorem 3. Setting ν to a number even slightly larger than 1 ensures that Algorithm 3 can converge to such a price region even in the extreme case where it starts with large amounts of over-demand. As shown in Figure 3, even in the MRVM domain where no linear clearing prices exist, the modifications of Algorithm 3 were sufficient for it to return a price vector with no predicted over-demand in *all* ML-powered clock rounds, in 100% of our instances (i.e., SATS seeds), while this number was almost 0% if we were to apply symmetrical GD as suggested by Theorem 3 by setting $\mu = \nu = 0$ in our Algorithm 3. Those results can be found in Appendix D.9.

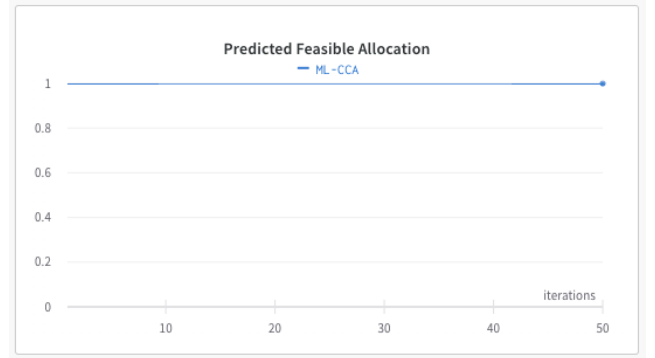


Figure 3: Fraction of instances in the MRVM domain where the price vector returned by Algorithm 3 was predicted to be feasible per iteration (starting after the Q_{init} -phase, i.e., in clock round 50). This fraction is constantly 100% across all rounds.

The following Example D.3 shows that even in cases where the constraint (9b) is mathematically irrelevant, the asymmetry of Algorithm 3 can still be very beneficial.

Example D.3. Let there be $m = 1$ item with capacity $c = 10$ and $n = 2$ bidders with value functions $v_1(x) = 6\mathbb{1}_{\{x \geq 6\}}$ and $v_2(x) = 3\mathbb{1}_{\{x \geq 1\}} + 2\mathbb{1}_{\{x \geq 5\}}$.³² First, note that similarly as in Example D.2 in this example also no LCP exists. For

³¹The actual MILP encoding we are using is based on the improved MILP-encoding of Weissteiner et al. (2023, Section 3.2 and Appendix F) and slightly modified to work with mMVNNs instead of classical MVNNs.

³²In the case of $m = 1$, we do not distinguish between the 1-dimensional vector x and the number x_1 (just as we write $p = p_1$ and $c = c_1$).

every $p < 0.5$ we have an over-demand of 1, and for every $p > 0.5$ we have an under-demand of at least 3 (3 for $p \in (0.5, 1]$, 9 for $p \in [1, 3]$ or 10 for $p \geq 3$). For $p = 0.5$, we have $\mathcal{X}_1^*(0.5) = \{6\}$ and $\mathcal{X}_2^*(0.5) = \{1, 5\}$ and thus the over-demand is either 1 or -3 . Thus, the price $p = 0.5$ is the unique minimizer of W and $p = 0.5$ also fulfills constraint (9b) (see Footnote 3). Therefore, $p = 0.5$ is both the unique solution to the constraint problem (9) and the unique solution to the unconstrained problem (37), which makes the two optimization problems (9) and (37) mathematically equivalent in this case.

However in practice, any GD-based algorithm will almost never be able to exactly compute $p = 0.5$. If we run the symmetric version of Algorithm 3 (see Remark D.3) we would either end up with a price p slightly below 0.5 or slightly above.³³ However, our asymmetric Algorithm 3 makes sure that we do not end up with a price p below 0.5, which would result in over-demand, but rather at a price p slightly above 0.5, which results at the feasible allocation where the first bidder gets 6 items and the second bidder gets 1 item. In this case, this would be the efficient allocation with a SCW of 9. If we had only asked demand queries for prices $p \in (0, 0.5)$ then the WDP would end up with a SCW of only 6 by allocating 6 items to first bidder and 0 items to the second bidder, since the first bidder would answer any of these demand queries with 6 items and the second bidder would answer them all with 5 items, which cannot be combined by the WDP for a capacity of $c = 10 < 11$.

D.7 Detailed Experimental Results

Detailed Efficiency Results In Tables 5 and 6, we provide the detailed efficiency results corresponding to Table 1 including 95%-bootstrapped CIs and p -values. Concretely, we now present triplets which show the lower bound of the bootstrapped 95%-CI, the mean, and the upper bound of the bootstrapped 95%-CI (e.g., (97.05, 97.87, 98.56) means that the lower bound of the bootstrapped 95%-CI is equal to 97.05, the mean is equal to 97.87, and the upper bound of the bootstrapped 95%-CI is equal to 98.56). Those bootstrapped CIs were created with the percentile method and 10.000 bootstrap-samples. It is noteworthy that in all domains other than SRVM (which is very easy, and can be solved by both mechanisms), our ML-CCA outperforms the CCA, both for the clock bids and the clock bids raised, and we can reject the null hypothesis of ML-CCA not outperforming the CCA in terms of efficiency at great confidence levels that, based on the domain, vary from less than 2% all the way to less than $7 \cdot 10^{-18}\% = 7e-20$.

³³In this case we would even end up more likely with a price below 0.5, because than we only have an over-demand of 1 rather than an under-demand of 3 items. In every step, where p^t is slightly below 0.5, GD will push up the price by only 1γ , while in the steps where p^t is slightly above 0.5, GD will push down the price by 3γ . The probability of exactly reaching $p = 0.5$ is zero if we initialize p^0 with a continuous distribution (as we do in Line 2). Also Line 31 is more likely to pick a price slightly below 0.5, since for every $\epsilon \in [0, 0.5)$, $W(0.5 + \epsilon, v) = W(0.5, v) + 3\epsilon$ and $W(0.5 - \epsilon, v) = W(0.5, v) + 1\epsilon$ (see Theorem 3). I.e., Line 31 would clearly prefer $p = 0.5 - \epsilon$ over $p = 0.5 + \epsilon$.

For GSVM and LSVM, for both practical bidding heuristics (clock bids and clock bids raised) the improvement of our ML-CCA over the CCA is highly significant with all p -values being below $2e-14$. For SRVM, the differences between the methods seem to be small, since both methods almost reach 100% efficiency. However, for clock bids raised this improvement is clearly statistically significant with a p -value of 0.0021%, while for clock bids there is actually no statistically significant difference. Note that for 7 out of all 8 practical settings (4 domains with two practical bidding heuristics) the p -value is below 1%. For clock bids raised all four domains have a p -value below 0.0021%. For MRVM, for all three bidding heuristics our ML-CCA is significantly better than the CCA with p -values 0.71%, $1.6e-5$ and $3e-4$.

In all three domains where LCPs exist, our ML-CCA found them statistically significantly more often than the CCA with p -values $3.2e-18$, $2.6e-8$ and 1.22% .

When reading the efficiency results in Tables 5 and 6 one should keep in mind that the CCA has generated over \$20 Billion in revenue for spectrum allocations between 2012 and 2014 alone (Ausubel and Baranov 2017). Since revenue is a lower bound for SCW, improving CCA's efficiency on average by 1% point would have improved the SCW by more than \$200 Million within this time range alone and CCA is still the most prominent practical mechanism for spectrum allocation.

Instead of efficiency $\frac{V(a)}{V(a^*)}$, one could also study the efficiency loss $\frac{V(a^*) - V(a)}{V(a^*)} = 1 - \frac{V(a)}{V(a^*)}$, which corresponds to the relative cost of deviating from the efficient allocation in terms of social welfare. For GSVM with clock bids only our ML-CCA can cut down the efficiency loss of the CCA by a factor 5.4 from 9.6% to 1.77%. Similarly for GSVM with clock bids raised our method can cut down the efficiency loss by a factor 5.9 from 6.41% to 1.07%.

Path Plots of Profit-Max Bids In Figure 4, we show the effect of adding up to $Q^{\text{P-Max}} = 100$ bids in the supplementary round of both our ML-CCA mechanism as well as the CCA using the profit-max heuristic. In GSVM, both mechanisms can reach 100% efficiency using those profit-max bids. However, ML-CCA's clock phase can do so after only 18 profit-max bids, while the CCA requires 44. In LSVM, for any number of profit max bids, our mechanism exhibits higher efficiency than the CCA, while with 100 profit-max bids we can reach 99.95% efficiency as opposed to 99.76% for the CCA. In SRVM, both mechanisms can reach over 99.99% efficiency using only 4 profit max bids. In MRVM, we can see that for almost any number of profit-max bids, our ML-CCA outperforms the CCA and the results are statistically significant on the 95% CI level.³⁴ Furthermore, for

³⁴Note that for example for 100 profit-max bids, the 95% CIs slightly overlap, while the corresponding p -value of the paired test is $0.03\% = 3e-4$. This suggest that a paired test would show the statistical significance of our ML-CCA having a higher average efficiency than the CCA for probably any number of profit-max bids. Note that a paired test is the correct statistical test for such situations, since both ML-CCA and the CCA were evaluated on the same 100 SATS-seeds.

MECHANISM	GSVM				LSVM			
	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA	(97.38, 98.23, 98.91)	(98.51, 98.93, 99.30)	(100.00, 100.00, 100.00)	56	(89.78, 91.64, 93.34)	(95.56, 96.39, 97.16)	(99.90, 99.95, 99.99)	26
CCA	(88.89, 90.40, 91.84)	(92.64, 93.59, 94.49)	(99.99, 100.00, 100.00)	3	(80.94, 82.56, 84.11)	(90.64, 91.60, 92.54)	(99.55, 99.76, 99.91)	0
p -VALUE	4.2E-18	6.5E-20	0.14	3.19E-18	1.8E-17	1.6E-14	0.0101	2.57E-8

Table 5: Detailed results for the GSVM and LSVM domains of ML-CCA vs CCA including the lower and upper 95%-bootstrapped CI bounds over a test set of 100 instances of the following metrics: E_{CLOCK} = efficiency in % for clock bids, E_{RAISE} = efficiency in % for raised clock bids, E_{PROFIT} = efficiency in % for raised clock bids and 100 profit-max demand queries, **CLEAR** = percentage of instances where linear clearing prices were found in the clock phase. Winners based on a paired t-test with $\alpha = 5\%$ are marked in grey. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA}} \leq \mu_{\text{CCA}}$ shows at which significance level we can reject the null hypothesis of CCA having a higher or equal average value in the corresponding metric than ML-CCA.

MECHANISM	SRVM				MRVM			
	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA	(99.48, 99.59, 99.68)	(99.92, 99.93, 99.95)	(100.00, 100.00, 100.00)	13	(92.79, 93.04, 93.28)	(93.09, 93.31, 93.52)	(93.46, 93.68, 93.89)	0
CCA	(99.52, 99.63, 99.73)	(99.75, 99.81, 99.86)	(100.00, 100.00, 100.00)	8	(91.87, 92.44, 92.86)	(92.25, 92.62, 92.96)	(92.84, 93.18, 93.48)	0
p -VALUE	0.78	2.1E-5	-	0.0122	7.1E-3	1.6E-5	3.0E-4	-

Table 6: Detailed results for the SRVM and MRVM domains of ML-CCA vs CCA including the lower and upper 95%-bootstrapped CI bounds over a test set of 100 instances of the following metrics: E_{CLOCK} = efficiency in % for clock bids, E_{RAISE} = efficiency in % for raised clock bids, E_{PROFIT} = efficiency in % for raised clock bids and 100 profit-max demand queries, **CLEAR** = percentage of instances where linear clearing prices were found in the clock phase. Winners based on a paired t-test with $\alpha = 5\%$ are marked in grey. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA}} \leq \mu_{\text{CCA}}$ shows at which significance level we can reject the null hypothesis of CCA having a higher or equal average value for the corresponding metric than ML-CCA.

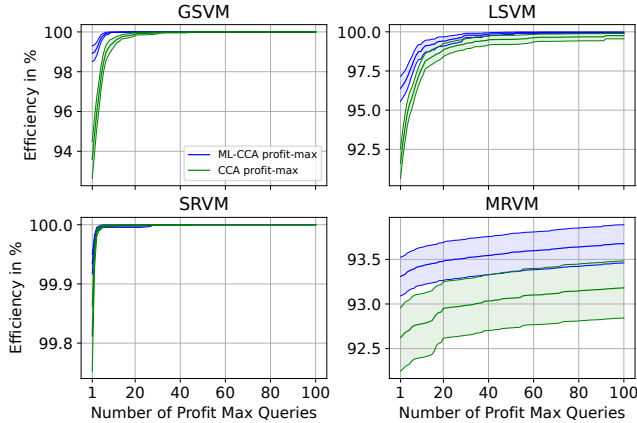


Figure 4: Efficiency of adding additionally $Q^{\text{P-Max}} = 100$ profit-max bids in SATS for ML-CCA and CCA after 100 clock bids and 100 raised clock bids. Averaged over 100 runs including a bootstrapped 95% CI.

MRVM, it is interesting to note that the CCA, even with 100 profit max bids per agent, cannot reach the clock bids raised efficiency of our ML-CCA (i.e., the efficiency with 0 profit max bids), while it needs 38 profit max bids to reach the efficiency that the clock phase of our ML-CCA achieves. In other words, the CCA requires up to an additional 138 value bids per agent to achieve the same efficiency that our ML-CCA can achieve using only 100 clock bids.

Path Plots of Clearing Error and Linear Item Prices In Figures 5 to 8, we present for all domains in SATS the path

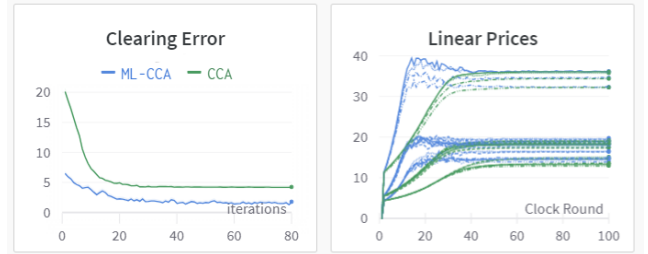


Figure 5: **Left:** CE of ML-CCA and CCA in GSVM per iteration (starting after the Q^{init} -phase, i.e., in clock round 20) averaged over 100 runs including the standard error. **Right:** linear item prices $p_j^r \in \mathbb{R}_{\geq 0}$ for $j = 1, \dots, 18$ defining the demand query for each clock round $r = 0, \dots, 100$ averaged over 100 runs.

plots for clock rounds $r = 0, \dots, 100$ of the (squared) clearing error (CE) defined as:

$$\sum_{j=1}^m \left(\left(\sum_{i=1}^n x_i^*(p^r)_j - c_j \right)^2 \right), \quad (84)$$

and the linear item prices $p_j^r \in \mathbb{R}_{\geq 0}$ for $j = 1, \dots, m$.

In Figure 5, we present the results for GSVM. We observe that, for any iteration after the initial $Q^{\text{init}} = 20$ clock rounds, i.e., any clock round $r = 20, \dots, 100$, the CE of our ML-CCA is smaller than the CE resulting from CCA. Specifically, ML-CCA has already at iteration 0 (clock round 20) a small CE (≈ 7) whilst CCA needs approximately 20 more iterations to reach a similar CE.



Figure 6: **Left:** CE of ML-CCA and CCA in LSVM per iteration (starting after the Q^{init} -phase, i.e., in clock round 20) averaged over 100 runs including the standard error. **Right:** linear item prices $p_j^r \in \mathbb{R}_{\geq 0}$ for $j = 1, \dots, 18$ defining the demand query for each clock round $r = 0, \dots, 100$ averaged over 100 runs.

Moreover, from the path plots of the prices we can distinguish the two phases of our proposed ML-CCA: the initial CCA phase with a predefined fixed price increment in case of over-demand for the first 20 clock rounds (or 50 in the case of MRVM), and the ML-powered demand query generation phase starting with the 21st (or 51st for MRVM) clock round. We observe that our approach, for the price of each item, is immediately searching in a local neighbourhood around the plateaued CCA price of that item, and tries to clear the market by slightly increasing and decreasing certain prices. Finally, we can see that CCA properly plateaus to final CCA prices around clock rounds 80 – 100, where no item is over-demanded anymore.

In Figure 6, we present the results for LSVM. We can see that for any iteration after the initial $Q^{\text{init}} = 20$ clock rounds, i.e., any clock round $r = 20, \dots, 100$, the CE of our ML-CCA is significantly smaller than the CE of CCA. Specifically, CCA requires 55 iterations (i.e., 75 clock rounds) to reach a CE comparable to what ML-CCA can achieve in the first iterations.³⁵ The price path plots in LSVM display a similar picture as in GSVM. We see that ML-CCA immediately identifies the correct region after the $Q^{\text{init}} = 20$ initial CCA prices and then tries to locally search by increasing and decreasing prices around the plateaued CCA prices, without at the same time sacrificing efficiency, as would be the case if we were to increase the price increment of CCA to reach that price area faster.

In Figure 7, we present the results for SRVM on a log scale. We can see again that, for any iteration after the initial $Q^{\text{init}} = 20$ clock rounds, i.e., any clock round $r = 20, \dots, 100$, the CE of our ML-CCA is smaller than that of CCA. Furthermore, we can see that in less than 20 iterations, our ML-CCA is able to drop the CE down to 1.1 (with 1 being a lower bound on the CE when clearing prices do not exist), while the CCA can never reach those numbers.

Recall, that in SRVM there are only three distinct items (i.e., $m = 3$) with quantities $c_1 = 6$, $c_2 = 14$, and $c_3 = 9$.

³⁵Note that by increasing the price increment for CCA, one could reduce the number of iterations until it achieves a low CE, but that could result in a significant drop in the efficiency of the auction, see Appendix D.8.

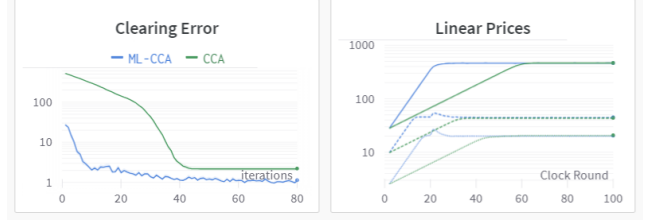


Figure 7: **Left:** CE of ML-CCA and CCA in SRVM per iteration (starting after the Q^{init} -phase, i.e., in clock round 20) averaged over 100 runs including the standard error. **Right:** linear item prices $p_j^r \in \mathbb{R}_{\geq 0}$ for $j = 1, \dots, 3$ defining the demand query for each clock round $r = 0, \dots, 100$ averaged over 100 runs. Both y-axes are on a log scale.

We again see the same behaviour of ML-CCA’s price discovery mechanism as in the other SATS domains: ML-CCA is able to immediately identify prices for the three items that are close to the final plateaued CCA prices.

In Figure 8, we present the results for MRVM. Interestingly, in this domain the CCA reaches a lower CE than our ML-CCA. However, please note that CE, as defined in Equation (84), penalizes all goods equally, which can be an uninformative metric in domains where the values of different goods vary significantly. This is most pronounced in the MRVM domain, where the prices of the most valuable items are more than 100 times larger than the prices of the least valuable goods, as shown in the right part Figure 8.

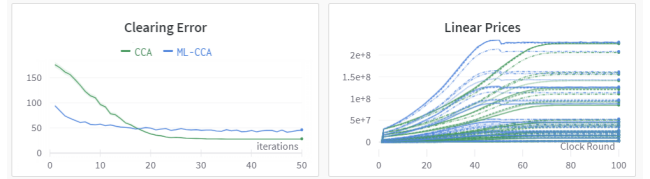


Figure 8: **Left:** CE of ML-CCA and CCA in MRVM per iteration (starting after the Q^{init} -phase, i.e., in clock round 50) averaged over 100 runs including the standard error. **Right:** linear item prices $p_j^r \in \mathbb{R}_{\geq 0}$ for $j = 1, \dots, 42$ defining the demand query for each clock round $r = 0, \dots, 100$ averaged over 100 runs.

Compute Time In Table 7 we report, for our choice of hyperparameters given in Table 3, the average time per round required in each domain to train the N mMVNNs of the bidders according to our Algorithm 1 and the time required to generate the next price vector, given the trained MVNNs, using our Algorithm 3. The sum of those two numbers is the average time per round required by our ML-CCA per domain to generate the next demand query, once the bidders have responded to the current one. For this experiment, we report average results over the same 100 instances as all other experiments in this paper.

Domain	Train Time	Price Gen. Time	Total
GSVM	4.49	1.81	6.30
LSVM	2.60	1.19	3.79
SRVM	1.26	0.41	1.67
MRVM	32.22	11.19	43.41

Table 7: Detailed results of the average time required (in minutes) for ML-CCA per round to train the mMVNNs of all bidders, and to generate the price vector of the next demand query, given the trained mMVNNs. Shown are average results over 100 instances.

D.8 Experimental Results for Reduced $Q^{\max} = 50$

In this section, we present all results of ML-CCA³⁶ and CCA for a reduced number of $Q^{\max} = 50$ demand queries instead of $Q^{\max} = 100$, which were presented in the main paper in Section 6. For ease of exposition we also include the CCA optimized for $Q^{\max} = 100$ as in Section 6. We additionally present 95% CIs and a paired t-test with $\alpha = 5\%$. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA}} \leq \mu_{\text{CCA}}$ shows at which significance level we can reject the null hypothesis of CCA with $Q^{\max} = 50$ having a higher or equal average value in the corresponding metric than the ML-CCA with $Q^{\max} = 50$. All results are presented in Tables 8 and 9.

Just like for $Q^{\max} = 100$, our ML-CCA outperforms the CCA considerably in all domains. Specifically, the efficiency improvement after the clock phase is over 7.3% points for GSVM, 8.6% points in LSVM and 3.4% points in MRVM, and all of those improvements are highly statistically significant, based on a paired t-test with $\alpha = 2.4\text{e}-10$. The two mechanisms are statistically tied in SRVM (other than for raised clock bids, where our ML-CCA again outperforms the CCA). As pointed out in Section 6, this domain is very easy, and both mechanisms can solve it even with 50 clock rounds). If we add the clock bids raised heuristic of the supplementary round to both mechanisms, the efficiency improvement of ML-CCA is 4.9% points in GSVM, 4.4% points in LSVM and 3.1% points in MRVM. Again all those improvements are highly significant with $\alpha = 2.9\text{e}-7$, this time also for SRVM and thus for all four domains. Finally, it is noteworthy that in the GSVM and LSVM domains, our ML-CCA with 50 clock rounds can achieve significantly higher efficiency with 50 clock rounds compared to the CCA with 100. These results even persist if we add the clock bids raised heuristic of the supplementary round to both mechanisms, which would induce up to an additional 100 value queries for each bidder in the CCA, and only 50 value queries for our ML-CCA.

D.9 Experimental Results for unconstrained W minimization

In this section, we empirically evaluate the importance of constraint (9b) in the NEXTPRICE-procedure by comparing

³⁶For ML-CCA, we used the same HPs for all domains as in Section 6, other than MRVM, where we also set $Q^{\text{init}} = 20$.

the efficiency results of ML-CCA with constraint (9b) (ML-CCA-C) and without it (ML-CCA-U).

Constraint (9b) ensures that the predicted induced demand should constitute a feasible allocation for every clock round of ML-CCA-C, as discussed in Section 4 and Appendix D.6.³⁷ On the other side, ML-CCA-U optimizes W without any constraint (i.e., by performing unconstrained classical GD on W as suggested by Lemma 2 and Theorem 3).³⁸

Those results are presented in Tables 10 and 11. In the GSVM, LSVM and SRVM domains, the results for the two approaches are almost identical (while constrained W minimization is statistically better in a few cases). A noteworthy difference is that for the GSVM and LSVM domains, the unconstrained W minimization clears the market in 5% and 3% less of the cases compared to the constrained version. Overall, as suggested in Section 4 and Appendix D.6, in domains where linear prices can achieve low clearing error (see Appendix D.7, Figures 5 to 7) minimizing W by performing classical GD on it without any additional constraints suffices for significant efficiency improvements compared to the CCA.

In the MRVM³⁹ domain however, we can see that the efficiency improvement of performing constrained W minimization compared to the unconstrained one is highly statistically significant across all three bidding heuristics with p -values ranging from $1.0\text{e}-8$ to $9.3\text{e}-6$. Without push bids, the efficiency improvement is approximately 1% point. This efficiency improvement also persists if we add the clock bids raised and profit max heuristics for the supplementary round, and is again highly statistically significant. Thus, we can conclude that in all cases, there is no disadvantage to using the constrained W minimization to generate the next price vector in ML-CCA. This clear improvements of ML-CCA-C over ML-CCA-U fits well our hypothesis stated in Remark 1 that constraint (9b) is especially important in cases where no LCPs exist (or when the market exhibits a high clearing error as defined in Equation (84), see Figures 5 to 8).

D.10 Experimental Results for using MVNNs on multiset domains

In this section, we experimentally evaluate the benefits of mMVNNs over MVNNs for multiset domains by comparing the efficiency results of ML-CCA with mMVNNs (ML-CCA-mMVNN) and MVNNs (ML-CCA-MVNN) as the neural network architecture. We only compare the two architectures for the two multiset domains, i.e., SRVM and MRVM, as in the other two domains, the two architectures are mathematically equivalent. It is important to note that,

³⁷ML-CCA-C is our default method, which we simply call ML-CCA outside of Appendix D.9. In particular, we used NEXTPRICE (Algorithm 3) with default hyper-parameters as discussed in Appendix D.6.

³⁸For ML-CCA-U we ignore constraint (9b), which corresponds to setting the hyper-parameters $\mu = \nu = 0$ in Algorithm 3 as described in Remark D.3.

³⁹For this test in MRVM, we used $Q^{\text{init}} = 70$ for both ML-CCA-C and ML-CCA-U.

MECHANISM	GSVM				LSVM			
	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA ($Q^{\max} = 50$)	(96.97, 97.84, 98.58)	(98.10, 98.59, 99.03)	(100.00, 100.00, 100.00)	51	(88.99, 90.81, 92.50)	(95.07, 95.91, 96.70)	(99.90, 99.95, 99.99)	17
CCA ($Q^{\max} = 50$)	(89.05, 90.51, 91.88)	(92.74, 93.70, 94.62)	(99.99, 100.00, 100.00)	1	(80.57, 82.21, 83.77)	(90.59, 91.52, 92.46)	(99.64, 99.80, 99.93)	0
CCA ($Q^{\max} = 100$)	(88.89, 90.40, 91.84)	(92.64, 93.59, 94.49)	(99.99, 100.00, 100.00)	3	(80.94, 82.56, 84.11)	(90.64, 91.60, 92.54)	(99.55, 99.76, 99.91)	0
p -VALUE ($Q^{\max} = 50$)	4.2E-17	4.2E-18	0.1433	7.0E-17	2.5E-16	3.1E-13	0.0145	9.1E-6

Table 8: Detailed results for the GSVM and LSVM domains of ML-CCA vs CCA for $Q^{\max} = 50$ including the lower and upper 95%-bootstrapped CI bounds over a test set of 100 instances of the following metrics: E_{CLOCK} = efficiency in % for clock bids, E_{RAISE} = efficiency in % for raised clock bids, E_{PROFIT} = efficiency in % for raised clock bids and 100 profit-max demand queries, CLEAR = percentage of instances where clearing prices were found in the clock phase. Winners for $Q^{\max} = 50$ based on a paired t-test with $\alpha = 5\%$ are marked in grey. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA}} \leq \mu_{\text{CCA}}$ shows at which significance level we can reject the null hypothesis of CCA with $Q^{\max} = 50$ having a higher or equal average value in the corresponding metric than ML-CCA with $Q^{\max} = 50$. Additionally, we also reprint the CCA ($Q^{\max} = 100$) results from Table 5 without marking statistical significance.

MECHANISM	SRVM				MRVM			
	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA ($Q^{\max} = 50$)	(99.47, 99.58, 99.67)	(99.92, 99.93, 99.95)	(100.00, 100.00, 100.00)	13	(91.63, 92.11, 92.50)	(91.94, 92.42, 92.81)	(92.73, 93.04, 93.34)	0
CCA ($Q^{\max} = 50$)	(99.36, 99.47, 99.59)	(99.64, 99.72, 99.79)	(100.00, 100.00, 100.00)	4	(87.70, 88.70, 89.61)	(88.52, 89.28, 90.02)	(89.57, 90.27, 90.92)	0
CCA ($Q^{\max} = 100$)	(99.52, 99.63, 99.73)	(99.75, 99.81, 99.86)	(100.00, 100.00, 100.00)	8	(91.87, 92.44, 92.86)	(92.25, 92.62, 92.96)	(92.84, 93.18, 93.48)	0
p -VALUE ($Q^{\max} = 50$)	0.088	2.9E-7	-	0.0011	2.4E-10	1.3E-12	1.6E-14	-

Table 9: Detailed results for the SRVM and MRVM domains of ML-CCA vs CCA for $Q^{\max} = 50$ including the lower and upper 95%-bootstrapped CI bounds over a test set of 100 instances of the following metrics: E_{CLOCK} = efficiency in % for clock bids, E_{RAISE} = efficiency in % for raised clock bids, E_{PROFIT} = efficiency in % for raised clock bids and 100 profit-max demand queries, CLEAR = percentage of instances where clearing prices were found in the clock phase. Winners for $Q^{\max} = 50$ based on a paired t-test with $\alpha = 5\%$ are marked in grey. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA}} \leq \mu_{\text{CCA}}$ shows at which significance level we can reject the null hypothesis of CCA with $Q^{\max} = 50$ having a higher or equal average value in the corresponding metric than ML-CCA with $Q^{\max} = 50$. In all domains (including MRVM) we used $Q^{\text{init}} = 20$ for this experiment. Additionally, we also reprint the CCA ($Q^{\max} = 100$) results from Table 6 without marking statistical significance.

in the case of MVNNs, even though the network has not incorporated the prior knowledge that some items are identical copies of each other, the price generation algorithm *does make use* of that prior information: in the case of MVNNs, our price generation algorithm appropriately calculates the total demand of each item as the aggregate demand of that item’s copies, and then sets the same price for all of those copies.

Those results are presented in Table 12. In the SRVM domain, the performance of the two architectures is almost identical in terms of both clearing potential and efficiency. The only statistically significant result is that mMVNNs slightly outperform MVNNs in terms of efficiency if the raised clock bids heuristic is used in the supplementary round. In the MRVM domain, (linear) clearing prices never exist, so no architecture ever clears the market. In terms of efficiency, the only statistically significant result in this domain is that MVNNs actually achieve 0.06% higher efficiency compared to mMVNNs after the clock phase. We have good reason to believe that the strong performance of MVNNs is because, as explained in the previous paragraph, the price generation algorithm *does make use* of the fact that some items are identical copies of each other, even though the MVNNs do not make use of that information. However, we are currently unable to verify this hypothesis, as the SATS simulator only supports demand queries where copies of the same item have identical prices.

MECHANISM	GSVM				LSVM			
	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA-C	(97.38, 98.23, 98.91)	(98.51, 98.93, 99.30)	(100.00, 100.00, 100.00)	56	(89.78, 91.64, 93.34)	(95.56, 96.39, 97.16)	(99.90, 99.95, 99.99)	26
ML-CCA-U	(97.05, 97.87, 98.56)	(98.22, 98.67, 99.07)	(100.00, 100.00, 100.00)	51	(89.84, 91.60, 93.28)	(95.33, 96.16, 96.94)	(99.90, 99.95, 99.99)	23
p -VALUE	0.0414	0.0965	-	0.0122	0.3428	0.0189	-	0.0416

Table 10: ML-CCA with constrained W (ML-CCA-C) and unconstrained W minimization (ML-CCA-U). Shown are averages including the lower and upper 95%-bootstrapped CI bounds over a test set of 100 synthetic CA instances for the GSVM and LSVM domains of the following metrics: efficiency in % for clock bids (E_{CLOCK}), raised clock bids (E_{RAISE}) and raised clock bids and 100 profit-max bids (E_{PROFIT}) and percentage of instances where linear clearing prices were found (CLEAR). Winners based on a paired t-test with $\alpha = 5\%$ are marked in grey. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA-C}} \leq \mu_{\text{ML-CCA-U}}$ shows at which significance level we can reject the null hypothesis of ML-CCA-U having a higher or equal average value in the corresponding metric than ML-CCA-C.

MECHANISM	SRVM				MRVM			
	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA-C	(99.48, 99.59, 99.68)	(99.92, 99.93, 99.95)	(100.00, 100.00, 100.00)	13	(92.90, 93.12, 93.34)	(93.03, 93.25, 93.47)	(93.46, 93.67, 93.88)	0
ML-CCA-U	(99.54, 99.63, 99.70)	(99.91, 99.93, 99.94)	(100.00, 100.00, 100.00)	13	(91.49, 92.06, 92.55)	(91.88, 92.31, 92.70)	(92.27, 92.70, 93.09)	0
p -VALUE	0.7793	0.0431	-	-	9.3E-6	2.7E-08	1.0E-8	-

Table 11: ML-CCA with constrained W (ML-CCA-C) and unconstrained W minimization (ML-CCA-U). Shown are averages including the lower and upper 95%-bootstrapped CI bounds over a test set of 100 synthetic CA instances for the SRVM and MRVM domains of the following metrics: efficiency in % for clock bids (E_{CLOCK}), raised clock bids (E_{RAISE}) and raised clock bids plus 100 profit-max bids (E_{PROFIT}) and percentage of instances where linear clearing prices were found (CLEAR). Winners based on a paired t-test with $\alpha = 5\%$ are marked in grey. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA-C}} \leq \mu_{\text{ML-CCA-U}}$ shows at which significance level we can reject the null hypothesis of ML-CCA-U having a higher or equal average value in the corresponding metric than ML-CCA-C.

MECHANISM	SRVM				MRVM			
	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR	E_{CLOCK}	E_{RAISE}	E_{PROFIT}	CLEAR
ML-CCA-mMVNN	(99.48, 99.59, 99.68)	(99.92, 99.93, 99.95)	(100.00, 100.00, 100.00)	13	(92.90, 93.12, 93.34)	(93.03, 93.25, 93.47)	(93.46, 93.67, 93.88)	0
ML-CCA-MVNN	(99.53, 99.63, 99.71)	(99.89, 99.91, 99.93)	(100.00, 100.00, 100.00)	13	(92.98, 93.18, 93.37)	(93.19, 93.39, 93.57)	(93.48, 93.69, 93.89)	0
p -VALUE	0.7401	0.0027	-	-	0.9700	0.8787	0.8182	-

Table 12: ML-CCA with mMVNNs (ML-CCA-mMVNN) and MVNNs (ML-CCA-MVNN) as the ML model architecture. Shown are averages including the lower and upper 95%-bootstrapped CI bounds over a test set of 100 synthetic CA instances for the SRVM and MRVM domains of the following metrics: efficiency in % for clock bids (E_{CLOCK}), raised clock bids (E_{RAISE}) and raised clock bids plus 100 profit-max bids (E_{PROFIT}) and percentage of instances where linear clearing prices were found (CLEAR). Winners based on a paired t-test with $\alpha = 5\%$ are marked in grey. The p -value for this pairwise t-test with $\mathcal{H}_0 : \mu_{\text{ML-CCA-mMVNN}} \leq \mu_{\text{ML-CCA-MVNN}}$ shows at which significance level we can reject the null hypothesis of ML-CCA-MVNN having a higher or equal average value in the corresponding metric than ML-CCA-mMVNN.