

No-Regret Learning and a Mechanism for Distributed Multiagent Planning

Jan-P. Calliess* Geoffrey J. Gordon

February 2008
CMU-ML-08-102



Report Documentation Page				Form Approved OMB No. 0704-0188	
Public reporting burden for the collection of information is estimated to average 1 hour per response, including the time for reviewing instructions, searching existing data sources, gathering and maintaining the data needed, and completing and reviewing the collection of information. Send comments regarding this burden estimate or any other aspect of this collection of information, including suggestions for reducing this burden, to Washington Headquarters Services, Directorate for Information Operations and Reports, 1215 Jefferson Davis Highway, Suite 1204, Arlington VA 22202-4302. Respondents should be aware that notwithstanding any other provision of law, no person shall be subject to a penalty for failing to comply with a collection of information if it does not display a currently valid OMB control number.					
1. REPORT DATE FEB 2008		2. REPORT TYPE		3. DATES COVERED 00-00-2008 to 00-00-2008	
4. TITLE AND SUBTITLE No-Regret Learning and a Mechanism for Distributed Multiagent Planning				5a. CONTRACT NUMBER	
				5b. GRANT NUMBER	
				5c. PROGRAM ELEMENT NUMBER	
6. AUTHOR(S)				5d. PROJECT NUMBER	
				5e. TASK NUMBER	
				5f. WORK UNIT NUMBER	
7. PERFORMING ORGANIZATION NAME(S) AND ADDRESS(ES) Carnegie Mellon University ,School of Computer Science,Pittsburgh,PA,15213				8. PERFORMING ORGANIZATION REPORT NUMBER	
9. SPONSORING/MONITORING AGENCY NAME(S) AND ADDRESS(ES)				10. SPONSOR/MONITOR'S ACRONYM(S)	
				11. SPONSOR/MONITOR'S REPORT NUMBER(S)	
12. DISTRIBUTION/AVAILABILITY STATEMENT Approved for public release; distribution unlimited					
13. SUPPLEMENTARY NOTES					
14. ABSTRACT We develop a novel mechanism for coordinated, distributed multiagent planning. We consider problems stated as a collection of single-agent planning problems coupled by common soft constraints on resource consumption. (Resources may be real or fictitious, the latter introduced as a tool for factoring the problem). A key idea is to recast the distributed planning problem as learning in a repeated game between the original agents and a newly introduced group of adversarial agents who influence prices for the resources. The adversarial agents benefit from arbitrage: that is, their incentive is to uncover violations of the resource usage constraints and, by selfishly adjusting prices, encourage the original agents to avoid plans that cause such violations. If all agents employ no-external-regret learning algorithms in the course of this repeated interaction, we are able to show that our mechanism can achieve design goals such as social optimality (efficiency) budget balance, and Nash-equilibrium convergence to within an error which approaches zero as the agents gain experience. In particular, the agents' average plans converge to a socially optimal solution for the original planning task. We present experiments in a simulated network routing domain demonstrating our method's ability to reliably generate sound plans.					
15. SUBJECT TERMS					
16. SECURITY CLASSIFICATION OF:			17. LIMITATION OF ABSTRACT Same as Report (SAR)	18. NUMBER OF PAGES 36	19a. NAME OF RESPONSIBLE PERSON
a. REPORT unclassified	b. ABSTRACT unclassified	c. THIS PAGE unclassified			

No-Regret Learning and a Mechanism for Distributed Multiagent Planning

Jan-P. Calliess* **Geoffrey J. Gordon[†]**

February 2008
CMU-ML-08-102

School of Computer Science
Carnegie Mellon University
Pittsburgh, PA 15213

*Machine Learning Dept., Carnegie Mellon University, Pittsburgh, PA, USA. The author is also a student at University of Karlsruhe, Germany.

[†]Machine Learning Dept., Carnegie Mellon University, Pittsburgh, PA, USA.

Keywords: Multiagent Planning, Multiagent Learning, No-Regret Learning, Online Convex Problems, Game Theory, Mechanism Design, Market-based Coordination.

Abstract

We develop a novel mechanism for coordinated, distributed multiagent planning. We consider problems stated as a collection of single-agent planning problems coupled by common soft constraints on resource consumption. (Resources may be real or fictitious, the latter introduced as a tool for factoring the problem). A key idea is to recast the distributed planning problem as learning in a repeated game between the original agents and a newly introduced group of *adversarial* agents who influence *prices* for the resources. The adversarial agents benefit from *arbitrage*: that is, their incentive is to uncover violations of the resource usage constraints and, by selfishly adjusting prices, encourage the original agents to avoid plans that cause such violations. If all agents employ no-external-regret learning algorithms in the course of this repeated interaction, we are able to show that our mechanism can achieve design goals such as social optimality (efficiency), budget balance, and Nash-equilibrium convergence to within an error which approaches zero as the agents gain experience. In particular, the agents' average plans converge to a socially optimal solution for the original planning task. We present experiments in a simulated network routing domain demonstrating our method's ability to reliably generate sound plans.

1 Introduction

In this work, we develop a novel, distributed multiagent planning mechanism. Our mechanism coordinates the different, individual goals of participating agents $P_1, \dots, P_k$ to achieve a globally desirable plan. While the agents could in principle compute the optimal global plan in a centralized manner, distributed approaches can improve robustness, fault tolerance, scalability (both in problem complexity and in the number of agents), and flexibility in changing environments [SD99].

We consider multiagent planning problems stated as $k \in \mathbb{N}$ single-agent convex optimization problems that are coupled by $n \in \mathbb{N}$ linear, soft constraints with positive coefficients. (By a *soft* constraint, we mean that violations are feasible, but are penalized by a convex loss function such as a hinge loss.) This representation includes, for example, network routing problems, in which each agent’s feasible region represents the set of paths from a source to a sink, its objective is to find a low-latency path, and a soft constraint represents the additional delay caused by congestion on a link used by multiple agents.

Since all coupling constraints (also referred to as *inter-agent constraints*) are soft, each agent’s feasible region does not depend on the actions of the other agents.¹ So, the agents could plan independently and be guaranteed that their joint actions would be feasible. This interaction is a *convex game* [SL07]: each player simultaneously selects her plan from a convex set, and if we hold all plans but one fixed, the remaining player’s loss is a convex function of her plan. By playing this convex game repeatedly, the players could learn about one another’s behavior and adjust their actions accordingly. With appropriate learning algorithms they could even ensure convergence to an equilibrium [BHLR07, GGMZ07]. Unfortunately, while distributed and robust, this naïve setup can lead to highly suboptimal global behavior: a selfish agent which can gain any benefit from using a congested link will do so, even if the resulting cost to other agents would far outweigh its own gain [Rou07].

To overcome this problem, a key idea of our approach is to introduce additional *adversarial* agents $A_1, \dots, A_n$, each of which can influence the cost of one of the resources by collecting usage fees. Like the original agents, the new agents are self-interested. But, collectively, they encourage the original agents to avoid excessive resource usage: we show below how to define their revenue functions so that they effectively perform *arbitrage*, allocating extra constraint-violation costs to each original agent in proportion to its responsibility for the violations.

The way we define the adversarial agents’ revenues and payments effectively *decouples* the individual planning problems: an original agent perceives the other original agents’ actions only through their effects on the choices of the adversarial agents. So, under our mechanism, the original agents never need to communicate with one another directly. Instead, they communicate with the adversarial agents to find out prices and declare demands for the resources they need to use. If an original agent never uses a particular resource, it never needs to send messages to the corresponding adversarial agent. This decoupling can greatly reduce communication requirements and increase robustness compared to the centralized solution: a central planner must communicate with every agent on every time step, and so constitutes a choke point for communication as well as

¹For a treatment on how our mechanism can be used in settings where inter-agent constraints are hard refer to Appendix E.

a single point of failure.

Because we have decoupled the agents from one another, each individual agent no longer needs to worry about the whole planning problem. Instead, it only has to solve a local online convex problem (OCP) [Gor99, Zin03] independently and selfishly. To solve its OCP, an agent could use any learning algorithm it desired; but, in this paper, we explore what happens if the agents use *no-regret* learning algorithms such as Greedy Projection [Zin03]. No-regret algorithms are a natural choice for agents in a multi-player game, since they provide performance guarantees that other types of algorithms do not. And, as we will show, if all agents commit to no-regret learning, several desirable properties result.

More specifically, if each agent's average per-iteration regret approaches zero, the agents will learn a globally optimal plan, in two senses: first, the *average per-iteration cost*, summed over all of the agents, will converge to the optimal social cost for the original planning problem. And second, the *average overall plan* of the original agents will converge to a socially optimal solution of the original planning problem.

These two results lead us to propose two different mechanism variants: in the **online** setup, motivated by the first result, learning takes place online in the classic sense; all agents choose and execute their plans and make and receive payments in every learning iteration. By contrast, in the **negotiation** setup, motivated by the second result, the agents learn and plan as usual, but only simulate the execution of their chosen joint plan. The simulated results (costs or rewards) from this proposed joint plan provide feedback for the learners. After a sufficient number of learning iterations, each agent averages together all of its proposed plans and executes the average plan. One can interpret either setup, but the negotiation setup in particular, as a very simple auction where the goods are resources. A plan which consumes a resource is effectively a bid for that resource; the resource prices are determined by the agents' learning behavior in response to the bids.

Just as with any mechanism, because our agents are selfish, we need to consider the impact of individual incentives. Focusing on the negotiation setup, we provide (mainly asymptotic) guarantees of Nash-equilibrium convergence of the overall learning outcome as well as classic mechanism design goals such as budget balance, individual rationality, and efficiency.

This document is a long version of material that appeared in the proceedings of the *7th International Conference of Autonomous Agents and Multiagent Systems (AAMAS 2008)* [CG08] and contains proofs that had to be omitted in the conference publication.

2 Preliminaries

In an *online convex program* [Gor99, Zin03], a possibly adversarial sequence $(\Gamma_{(t)})_{t \in \mathbb{N}}$ of convex cost functions is revealed step by step. (Equivalently, one could substitute concave reward functions.) At each step t , the OCP algorithm must choose a play $\mathbf{x}_{(t)}$ from its feasible region F while only knowing the *past* cost functions $\Gamma_{(q)}$ and choices $\mathbf{x}_{(q)}$ ($q \leq t - 1$). After the choice is made, the current cost function $\Gamma_{(t)}$ is revealed, and the algorithm pays $\Gamma_{(t)}(\mathbf{x}_{(t)})$.

To measure the performance of an OCP algorithm, we can compare its accumulated cost up through step T to an estimate of the best cost attainable against the sequence $(\Gamma_{(t)})_{t=1 \dots T}$. Here, we will estimate the best attainable cost as the cost of the best constant play $\mathbf{s}_{(T)} \in F$, chosen with

knowledge of $\Gamma_{(1)} \dots \Gamma_{(T)}$. This choice leads to a measure called *external regret* or just *regret*: $R(T) = \sum_{t=1}^T \Gamma_{(t)}(\mathbf{x}_{(t)}) - \sum_{t=1}^T \Gamma_{(t)}(\mathbf{s}_{(T)})$. An algorithm is *no-(external)-regret* iff it guarantees that $R(T)$ grows slower than $O(T)$, i.e., $R(T) \leq \Delta(T) \in o(T)$. $\Delta(T)$ is a *regret bound*. (The term *no-regret* is motivated by the fact that the limiting average regret is no more than zero, $\limsup_{T \rightarrow \infty} R(T)/T \leq 0$.)

We define the *convex conjugate* or *dual* [BV04] of a function $\Gamma(\mathbf{x})$ to be $\Gamma^*(\mathbf{y}) = \sup_{\mathbf{x} \in \text{dom } \Gamma} [\langle \mathbf{x}, \mathbf{y} \rangle - \Gamma(\mathbf{x})]$. The conjugate function Γ^* is always closed and convex, and if Γ is closed and convex, then $\Gamma^{**} = \Gamma$ pointwise.

2.1 Model and Notation

We wish to model interaction among k *player* agents, $P_1 \dots P_k$. We represent player P_i 's individual planning problem as a convex program: choose a vector $\mathbf{p}_i$ from a compact, convex feasible set F_{P_i} to minimize the *intrinsic* cost $\langle \mathbf{c}_i, \mathbf{p}_i \rangle$. (In addition to the intrinsic cost, player P_i will attempt to minimize cost terms which arise from interactions with other agents; we will define these extra terms below and add them to P_i 's objective.) We assume that the intrinsic cost vector $\mathbf{c}_i$ and feasible region F_{P_i} are private information for P_i —that is, P_i may choose to inform other players about them, but may also choose to be silent or to lie.

We assume each feasible set F_{P_i} is a subset of some finite-dimensional $\mathbb{R}$ -Hilbert space V_{P_i} with standard scalar product $\langle \cdot, \cdot \rangle_{V_{P_i}}$. We write $V = V_{P_1} \times \dots \times V_{P_k}$ and $F_P = F_{P_1} \times \dots \times F_{P_k}$ for the *overall* planning space and feasible region, and $\mathbf{p} = (\mathbf{p}_1; \dots; \mathbf{p}_k)$ and $\mathbf{c} = (\mathbf{c}_1; \dots; \mathbf{c}_k)$ for the overall plan and combined objective. (We use $;$ to denote vertical stacking of vectors, consistent with Matlab usage.) And, for any $\mathbf{c}, \mathbf{p} \in V$, we write $\langle \mathbf{c}, \mathbf{p} \rangle_V = \sum_i \langle \mathbf{c}_i, \mathbf{p}_i \rangle_{V_{P_i}}$ for our scalar product on V . (We omit subscripts on $\langle \cdot, \cdot \rangle$ when doing so will not cause confusion.) For convenience, we assume each feasible set only contains vectors with nonnegative components; we can achieve this property by changing coordinates if necessary.

We model the coupling among players by a set of n soft linear constraints, which we interpret as resource consumption constraints. That is, we assume that there are vectors $\mathbf{l}_{ji} \geq 0$ such that the consumption of resource j by plan $\mathbf{p}_i \in V_{P_i}$ is $\langle \mathbf{l}_{ji}, \mathbf{p}_i \rangle$. (So, the total consumption of resource j is $\langle \mathbf{l}_j, \mathbf{p} \rangle$, where $\mathbf{l}_j = (\mathbf{l}_{j1}; \dots; \mathbf{l}_{jk})$.) And, we assume that there are monotone increasing, continuous, convex penalty functions $\beta_j(\nu)$ and scalars $y_j \geq 0$ so that the overall cost due to consumption of resource j in plan $\mathbf{p}$ is $\beta_j(\langle \mathbf{l}_j, \mathbf{p} \rangle - y_j)$. In keeping with the interpretation of β_j as enforcing a soft constraint, we assume $\beta_j(\nu) = 0$ for all $\nu \leq 0$. (For example, $\beta_j(\nu)$ could be the hinge loss function $\max\{0, \nu\}$.) We define

$$\nu_j(\mathbf{p}) = \langle \mathbf{l}_j, \mathbf{p} \rangle - y_j \quad (1)$$

to be the magnitude of violation of the j th soft resource constraint by plan $\mathbf{p}$.² Because the resource constraints are soft, no player can make another player's chosen action infeasible; conflicts result only in high cost rather than infeasibility.

The function β_j describes the *overall* cost to all players of usage of resource j . We will assume that, in the absence of any external coordination mechanism, the cost to player i due to resource

²To enforce a hard constraint, we could choose a sufficiently small margin $\epsilon > 0$, replace y_j by $y_j - \epsilon$, and set $\beta_j(\nu) = \max\{0, \nu/\epsilon\}$.

j is given by some function $\beta_{ji}(\mathbf{p})$ with $\sum_i \beta_{ji}(\mathbf{p}) = \beta_j(\nu_j(\mathbf{p}))$. We will call β_{ji} the *natural cost* or *cost in nature* of P_i 's resource usage. A typical choice for β_{ji} is proportional to player i 's consumption of resource j :

$$\beta_{ji}(\mathbf{p}) = \begin{cases} 0 & , \text{ if } \langle \mathbf{l}_j, \mathbf{p} \rangle = 0 \\ \beta_j(\nu_j(\mathbf{p})) \langle \mathbf{l}_{ji}, \mathbf{p}_i \rangle / \langle \mathbf{l}_j, \mathbf{p} \rangle & , \text{ otherwise} \end{cases} \quad (2)$$

So, including both her intrinsic costs and the natural costs of resource usage, player i 's objective is

$$\omega_{P_i}(\mathbf{p}) = \langle \mathbf{c}_i, \mathbf{p}_i \rangle + \sum_j \beta_{ji}(\mathbf{p}) .$$

We will write $\omega(\mathbf{p}) = \sum_i \omega_{P_i}(\mathbf{p})$ for the social cost; as stated above, our goal is to coordinate the player agents to minimize $\omega(\mathbf{p})$. (With this notation, several facts mentioned above should now be obvious: for example, since the individual objectives ω_{P_i} depend on the entire joint plan $\mathbf{p}$, the players cannot simply plan in isolation. Nor do we want the players to compute and follow an equilibrium: using the above choice for β_{ji} (which results in a setting similar to so-called *nonatomic congestion games*), there are simple sequences of examples showing that the penalty for following an equilibrium (called the *price of anarchy*) can be arbitrarily large [Rou07].)

3 Problem transformation

In this section, by dualizing our soft resource constraints, we decouple the problem of finding a socially optimal plan. The result is a *saddle-point* or *minimax* problem whose variables are the original plan vector $\mathbf{p}$ along with new dual variables $\mathbf{a}$, defined below. By associating the new variables $\mathbf{a}$ with additional agents, called *adversarial agents*, we arrive at a convex game with comparatively-sparse interactions. Based on this game, we introduce our proposed learning-based mechanism.

3.1 Introduction of adversarial agents

Write $F_{A_j} = \text{dom } \beta_j^*$. Since we have assumed that β_j is continuous and convex, we know that $\beta_j^{**} = \beta_j$ pointwise, that is, $\beta_j(\nu) = \sup_{a^j \in F_{A_j}} [a^j \nu - \beta_j^*(a^j)]$ for all $\nu \in \mathbb{R}$. Since β_j^* will become part of the objective function for the adversarial agents, and since many online learning algorithms require compact domains, we will assume that $F_{A_j} = [0, u^j]$ for some scalar u^j . (For example, this assumption is satisfied if the slope of β_j is upper bounded by u_j , and achieves its upper bound. The lower bound of zero follows from our previous assumptions that β_j is monotone and $\beta_j(\nu) = 0$ for $\nu \leq 0$.) We will also assume that β_j^* is continuous on its domain.

We define $\Omega_{A_j} : V \times F_{A_j} \rightarrow \mathbb{R}$ as

$$\Omega_{A_j}(\mathbf{p}, a^j) = a^j \nu_j(\mathbf{p}) - \beta_j^*(a^j) . \quad (3)$$

And, writing $\mathbf{a} = (a^1; \dots; a^n) \in F_A = F_{A_1} \times \dots \times F_{A_n}$, we define³

³Note, in the AAMAS version [CG08] the explicit mention of the restriction of Ω to feasible plans was omitted. However, whenever we considered saddle-points we also considered them with respect to this restricted function $\Omega : F_P \times F_A \rightarrow \mathbb{R}$.

$$\Omega : \begin{cases} F_P \times F_A \rightarrow \mathbb{R} \\ (\mathbf{p}, \mathbf{a}) \mapsto \langle \mathbf{c}, \mathbf{p} \rangle + \sum_{j=1}^n \Omega_{A_j}(\mathbf{p}, a^j) \end{cases} \quad (4)$$

(note the inclusion of the intrinsic cost $\langle \mathbf{c}, \mathbf{p} \rangle$). Because of the duality identity mentioned above, along with our assumption about $\text{dom } \beta_j^*$, we know that for all plans $\mathbf{p}$, $\sup_{a^j \in F_{A_j}} \Omega_{A_j}(\mathbf{p}, a^j) = \beta_j(\nu_j(\mathbf{p}))$, and so

$$\omega(\mathbf{p}) = \max_{\mathbf{a} \in F_A} \Omega(\mathbf{p}, \mathbf{a}), \forall \mathbf{p} \in F_P. \quad (5)$$

Note that we have replaced $\sup$ by $\max$ in Eq. 5: since $\Omega(\mathbf{p}, \cdot)$ is a closed concave function, it achieves its supremum on a compact domain such as F_A .

Remark 3.1. Note, Eq. 5 establishes the connection between saddle-points of Ω and overall player plans that collectively minimize total cost in nature: If $(\tilde{\mathbf{p}}, \tilde{\mathbf{a}})$ is a saddle-point then we have

$$\tilde{\mathbf{p}} = \underset{\mathbf{p} \in F_P}{\operatorname{argmin}} \max_{\mathbf{a} \in F_A} \Omega(\mathbf{p}, \mathbf{a}) = \underset{\mathbf{p} \in F_P}{\operatorname{argmin}} \omega(\mathbf{p}). \quad (6)$$

(Cf. [Roc70]).

Now, as promised, we can introduce the *adversarial agents*: the adversarial agent A_j controls the parameter $a^j \in F_{A_j}$, and tries to maximize its revenue $\Omega_{A_j}(\mathbf{p}, a^j) - \beta_j(\nu_j(\mathbf{p}))$. Note that $\beta_j(\nu_j(\mathbf{p}))$ does not depend on a^j , and so does not affect the choice of a^j once $\mathbf{p}$ is fixed.

To give A_j this revenue, we will have player P_i pay adversary A_j the amount $a^j \langle \mathbf{l}_{ji}, \mathbf{p}_i \rangle - \beta_{ji}(\mathbf{p}) - d_{ji} r_j(a^j)$. Here the *remainder function* r_j is defined as $r_j(a^j) = a^j y_j + \beta_j^*(a^j)$; the nonnegative weights d_{ji} are responsible for dividing up the remainder among all player agents, so we require $\sum_i d_{ji} = 1$ for each j . Given these definitions, it is easy to check that the sum of all payments to A_j is indeed $\Omega_{A_j}(\mathbf{p}, a^j) - \beta_j(\nu_j(\mathbf{p}))$ as claimed.

We can interpret the above payments as follows: A_j sets the per-unit price a^j for consumption of resource j . P_i pays A_j according to consumption, $a^j \langle \mathbf{l}_{ji}, \mathbf{p}_i \rangle$, and is reimbursed for her share of the actual resource cost, $\beta_{ji}(\mathbf{p})$. For the privilege of setting the per-unit price, A_j pays a fee $r_j(a^j)$; this fee is distributed back to the player agents according to weights d_{ji} . (We show in Appendix D that the fee is always nonnegative.) Since the entire revenue for the agents A_j arises from payments by the player agents, we can think of A_j as opponents for the players—this is qualitatively true even though our game has many players and even though the player agent payoff functions contain terms that do not involve $\mathbf{a}$.

Including payments to adversaries, P_i 's cost becomes

$$\Omega_{P_i}(\mathbf{p}_i, \mathbf{a}) = \langle \mathbf{c}_i, \mathbf{p}_i \rangle + \sum_j (a^j \langle \mathbf{l}_{ji}, \mathbf{p}_i \rangle - d_{ji} r_j(a^j)). \quad (7)$$

By the above construction, we have achieved several important properties:

- First, as promised, P_i 's cost does not depend on any components of $\mathbf{p}$ other than $\mathbf{p}_i$, and A_j 's revenue does not depend on any components of $\mathbf{a}$ other than a^j . So, given an adversarial play $\mathbf{a}$, each player could plan by independently optimizing $\Omega_{P_i}(\mathbf{p}_i, \mathbf{a})$. Similarly, given a plan $\mathbf{p}$, each adversary could separately optimize $\Omega_{A_j}(\mathbf{p}, a_j)$. (A_j can ignore the term $\beta_j(\nu_j(\mathbf{p}))$ if $\mathbf{p}$ is fixed.) So, the players' optimization problems are *decoupled* given $\mathbf{a}$, and the adversaries' optimization problems are *decoupled* given $\mathbf{p}$.

- Second, if an adversarial agent plays optimally, her revenue will be exactly zero, since $\max_{a^j \in [0, u^j]} \Omega_{A_j}(\mathbf{p}, a^j) = \beta_j(\nu_j(\mathbf{p}))$. (Suboptimal play will lead to a negative revenue, i.e., a loss or cost.)
- Third, the total cost to all player agents is

$$\sum_i \Omega_{P_i}(\mathbf{p}_i, \mathbf{a}) = \Omega(\mathbf{p}, \mathbf{a}) . \quad (8)$$

On the other hand, the total revenue to all adversaries is $\Omega_A(\mathbf{p}, \mathbf{a}) = \Omega(\mathbf{p}, \mathbf{a}) - \langle \mathbf{c}, \mathbf{p} \rangle - \sum_j \beta_j(\nu_j(\mathbf{p}))$. If the adversaries each play optimally, then $\Omega_A(\mathbf{p}, \mathbf{a})$ will be zero, so we will have

$$\Omega(\mathbf{p}, \mathbf{a}) = \langle \mathbf{c}, \mathbf{p} \rangle + \sum_j \beta_j(\nu_j(\mathbf{p})) = \omega(\mathbf{p}) . \quad (9)$$

Combining Eqs. 8 and 9, we find that if the adversaries play optimally, the total cost to the player agents is $\omega(\mathbf{p})$, just as it was in our original planning problem.

- Finally, since $\Omega(\mathbf{p}, \mathbf{a})$ is a continuous saddle-function on a compact domain [Roc70], it must have a saddle-point $(\tilde{\mathbf{p}}, \tilde{\mathbf{a}})$. (By definition, a saddle-point is a point $(\tilde{\mathbf{p}}, \tilde{\mathbf{a}})$ such that $\Omega(\mathbf{p}, \tilde{\mathbf{a}}) \geq \Omega(\tilde{\mathbf{p}}, \tilde{\mathbf{a}}) \geq \Omega(\tilde{\mathbf{p}}, \mathbf{a})$ for all $\mathbf{p} \in F_P$ and $\mathbf{a} \in F_A$.) By the decoupling arguments above, we must have that $\tilde{\mathbf{p}}_i \in \arg \min_{\mathbf{p}_i \in F_{P_i}} \Omega_{P_i}(\mathbf{p}_i, \tilde{\mathbf{a}})$ for each i , and that $\tilde{\mathbf{a}}_j \in \arg \max_{\mathbf{a}_j \in F_{A_j}} \Omega_{A_j}(\tilde{\mathbf{p}}, \mathbf{a}_j)$ for each j . (The latter is true since Ω and Ω_A differ only by terms that do not depend on $\mathbf{a}$.)

3.2 Planning as learning in a repeated game

If we consider our planning problem as a game among all of the agents P_i and A_j , we have just shown that there exists a Nash equilibrium in pure strategies, and that in any Nash equilibrium, the player plan $\tilde{\mathbf{p}}$ must minimize $\omega(\tilde{\mathbf{p}})$ and therefore be socially optimal. To allow the agents to find such an equilibrium, we now cast the planning problem as learning in a repeated game. We will show that, if each agent employs a no-regret learning algorithm, the agents as a whole will converge to a socially optimal plan, both in the sense that the average joint plan converges and in the sense that the average social cost converges. (This result, while similar to well-known results about convergence of no-regret algorithms to minimax equilibrium [FS96], does not follow from these results, in part because our game is not constant-sum.) Note that, from the individual agent's perspective, playing in the repeated game is an OCP, and so using a no-regret learner would be a reasonable choice; we explore the effect of this choice in more detail below in Sec. 4.

The repeated game is played between the k players and the n adversarial agents. Based on their local histories of past observations, in each round t , each player P_i chooses a current pure strategy $\mathbf{p}_{i(t)} \in F_{P_i}$, and simultaneously, each adversary A_j chooses a current resource price $a_{(t)}^j \in F_A$. We write $\mathbf{p}_{(t)} = (\mathbf{p}_{1(t)}; \dots; \mathbf{p}_{k(t)})$ and $\mathbf{a}_{(t)} = (a_{(t)}^1; \dots; a_{(t)}^n)$ for the joint actions of all players and adversaries, respectively.

After choosing $\mathbf{p}_{(t)}$ and $\mathbf{a}_{(t)}$, the players send their current resource consumptions $\langle \mathbf{l}_{ji}, \mathbf{p}_{i(t)} \rangle$ to the adversaries, and the adversaries send their current prices to the players. In the online model, P_i

observes $\beta_{ji}(\mathbf{p}_{(t)})$ and sends it to A_j as well; in the negotiation model, we assume that β_{ji} is of the form given in Eq. 2, so that A_{ji} can compute $\beta_{ji}(\mathbf{p}_{(t)})$. The above information allows each P_i to compute its current cost function $\Omega_{P_i(t)}(\cdot) = \Omega_{P_i}(\cdot, \mathbf{a}_{(t)})$ and its cost $\Omega_{P_i(t)}(\mathbf{p}_{i(t)})$. It also allows each A_j to compute $\Omega_{A_j(t)}(\cdot) = \Omega_{A_j}(\mathbf{p}_{(t)}, \cdot)$ as well as $\beta_{j(t)} = \beta_j(\nu_j(\mathbf{p}_{(t)}))$, and thus, its total revenue $\Omega_{A_j(t)}(a_{j(t)}^j) - \beta_{j(t)}$. (In fact, A_j may avoid computing or storing $\beta_{j(t)}$ if desired, since it does not influence that term directly.) In Sec. 3.3 below, we discuss how to implement the necessary communication efficiently.

Each player P_i then adds observation $\mathbf{p}_{(t)}$, $\Omega_{P_i(t)}(\cdot)$ to her local history, and each adversary A_j adds $\mathbf{a}_{(t)}$, $\Omega_{A_j(t)}(\cdot)$ to her local history. Finally, the system enters iteration $t + 1$, and the process repeats.

3.2.1 Game between two synthesized agents

For analysis, it will help to construe our setup as a fictitious game between a *synthesized player agent*, P , and a *synthesized adversarial agent*, A . When each component agent P_i plays $\mathbf{p}_{i(t)}$ and each component agent A_j plays $a_{j(t)}^j$, then we imagine P to play $\mathbf{p}_{(t)}$ and A to play $\mathbf{a}_{(t)}$. Accordingly, we understand P to incur cost $\Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)})$, and A to have revenue $\Omega_A(\mathbf{p}_{(t)}, \mathbf{a}_{(t)}) - \sum_j \beta_j(\nu_j(\mathbf{p}_{(t)}))$ in round t . These synthesized agents are merely theoretical notions serving to simplify our reasoning; in practice there would never be a single agent controlling all players.

Using these synthesized agents, we will prove two results: first, immediately below, we show that if the individual agents use no-regret algorithms, then the synthesized agents also achieve no regret. And second, in Sec. 3.2.2, we show that if the synthesized agents achieve no regret, then they will converge to an equilibrium of the game, in the two senses mentioned above.

Lemma 3.2. *If each individual agent P_i achieves regret bound $\Delta_{P_i}(T)$, then the synthesized player agent P achieves regret bound $\Delta_P(T) := \sum_i \Delta_{P_i}(T)$. So, if $\Delta_{P_i}(T) \in o(T)$ for all i , then $\Delta_P(T) \in o(T)$.*

Proof. By definition, the regret $R_P(T)$ for agent P is

$$R_P(T) = \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)}) - \min_{\mathbf{p}} \sum_{t=1}^T \Omega(\mathbf{p}, \mathbf{a}_{(t)})$$

and the regret for P_i is $R_{P_i}(T) \leq \Delta_{P_i}(T)$:

$$R_{P_i}(T) = \sum_{t=1}^T \Omega_{P_i(t)}(\mathbf{p}_{i(t)}) - \min_{\mathbf{p}_i} \sum_{t=1}^T \Omega_{P_i(t)}(\mathbf{p}_i)$$

Owing to the decoupling effect of the adversary we have $\Omega(\mathbf{p}, \mathbf{a}_{(t)}) = \sum_{i=1}^k \Omega_{P_i(t)}(\mathbf{p}_i)$. So, we can

expand $R_P(T)$ as

$$\begin{aligned}
& \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)}) - \min_{\mathbf{p} \in F_P} \sum_{t=1}^T \Omega(\mathbf{p}, \mathbf{a}_{(t)}) \\
&= \sum_{t=1}^T \sum_{i=1}^k \Omega_{P_{i(t)}}(\mathbf{p}_{i(t)}) - \min_{\mathbf{p} \in F_P} \sum_{t=1}^T \sum_{i=1}^k \Omega_{P_{i(t)}}(\mathbf{p}_i) \\
&= \sum_{i=1}^k \sum_{t=1}^T \Omega_{P_{i(t)}}(\mathbf{p}_{i(t)}) - \sum_{i=1}^k \min_{\mathbf{p}_i \in F_{P_i}} \sum_{t=1}^T \Omega_{P_{i(t)}}(\mathbf{p}_i) \\
&= \sum_{i=1}^k \left(\sum_{t=1}^T \Omega_{P_{i(t)}}(\mathbf{p}_{i(t)}) - \min_{\mathbf{p}_i \in F_{P_i}} \sum_{t=1}^T \Omega_{P_{i(t)}}(\mathbf{p}_i) \right) \\
&= \sum_{i=1}^k R_{P_i}(T) \leq \sum_{i=1}^k \Delta_{P_i}(T) .
\end{aligned}$$

So, $R_P(T) \in o(T)$ as desired. $\square$

Analogously, for the adversarial agent A we have the following lemma. The proof is very similar, and is therefore omitted.

Lemma 3.3. *If each adversarial agent A_j achieves regret bound $\Delta_{A_j}(T)$, then the synthesized agent A achieves regret bound $\Delta_A(T) := \sum_j \Delta_{A_j}(T)$. So, if $\Delta_{A_j}(T) \in o(T)$ for all j , then $\Delta_A(T) \in o(T)$.*

3.2.2 Social optimality

In this section, we investigate the behavior of the averaged strategies $\bar{\mathbf{p}}_{[T]} := \frac{1}{T} \sum_{t=1}^T \mathbf{p}_{(t)}$ and $\bar{\mathbf{a}}_{[T]} := \frac{1}{T} \sum_{t=1}^T \mathbf{a}_{(t)}$, as well as the averaged costs $\frac{1}{T} \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)})$. (Recall that the negotiation version of our mechanism outputs the averaged strategies, while the online version of our mechanism incurs the averaged costs.)

Starting with the averaged strategies, we show that if all players achieve no regret, we can guarantee convergence of $(\bar{\mathbf{p}}_{[T]}, \bar{\mathbf{a}}_{[T]})$ to a set $K_P \times K_A$ of saddle-points of Ω (where $K_P \subset F_P, K_A \subset F_A$). (While the sequences $\bar{\mathbf{p}}_{[T]}$ and $\bar{\mathbf{a}}_{[T]}$ may not converge, the distance of $\bar{\mathbf{p}}_{[T]}$ from K_P and the distance of $\bar{\mathbf{a}}_{[T]}$ from K_A will approach zero; and, every cluster point of the sequence $(\bar{\mathbf{p}}_{[T]}, \bar{\mathbf{a}}_{[T]})$ will be in $K_P \times K_A$. Due to the convexity and compactness of the feasible sets, each average strategy and each cluster point will be feasible.)

Theorem 3.4. *Let $\bar{\mathbf{p}}_{[T]} = \frac{1}{T} \sum_{t=1}^T \mathbf{p}_{(t)}$, $\bar{\mathbf{a}}_{[T]} = \frac{1}{T} \sum_{t=1}^T \mathbf{a}_{(t)}$ be the averaged, pure strategies of synthesized player P and adversary A , respectively. If P, A each suffer sublinear external regret, then as $T \rightarrow \infty$, $(\bar{\mathbf{p}}_{[T]}, \bar{\mathbf{a}}_{[T]})$ converges to a (bounded) subset $K_P \times K_A$ of saddle-points of the player cost function $\Omega : F_P \times F_A \rightarrow \mathbb{R}$.*

Proof. Refer to Appendix C. $\square$

Since Ω is continuous on its domain, Thm. 3.4 lets us conclude that the outcome of negotiation is a plan which approximately minimizes total player cost in nature: if we choose T sufficiently large, then $(\bar{\mathbf{p}}_{[T]}, \bar{\mathbf{a}}_{[T]})$ must be close to a saddle-point, and so the costs must be close to the costs of a saddle-point. (To determine how large we need to choose T , we can look at the regret bounds of the learning algorithms of the individual agents.) As expressed in Rem. 3.1, the player part $\tilde{\mathbf{p}}$ of a saddle-point incurs minimal total player cost in nature and, since the adversarial agents learn to set prices as nature would, $\tilde{\mathbf{p}}$ is socially optimal (w.r.t. to the player agents) with respect to both nature and the mechanism.

If we choose to run our mechanism in online mode, we also need bounds on the average incurred cost. Since $\omega(\mathbf{p}) = \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a})$, the following theorem tells us that the average social cost approaches the optimal social cost in the long run.

Theorem 3.5. *If P and A suffer sublinear external regret, then as $T \rightarrow \infty$,*

$$\frac{1}{T} \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)}) \rightarrow \min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) .$$

Proof. Refer to Appendix C. □

3.3 Communication costs

So far we have assumed that all player agents broadcast their resource usages (and possibly their natural costs) to all adversarial agents, and all adversarial agents broadcast their prices to all player agents. With this assumption, on every time step, each player sends one broadcast of size $O(n)$ (her resource usages) and receives n messages of size $O(1)$ (the resource prices), while each adversary sends one broadcast of size $O(1)$ and receives k messages of size $O(n)$, for a total of $n + k$ broadcasts per step, and a total incoming bandwidth of no more than $O(nk)$ at each agent.⁴ Even under this simple assumption, the cost is somewhat better than a centralized planner, which would have to receive k much larger messages describing each player’s detailed optimization problem, and send k much larger messages describing each player’s optimal plan.

However, by exploiting *locality*, we can reduce bandwidth even further: in many problems we can guarantee *a priori* that player P_i will never use resource j , and in this case, we never need to transmit A_j ’s price a^j to P_i . Similarly, if P_i decides not to use resource j on a given trial, we never need to transmit $\langle \mathbf{l}_{ji}, \mathbf{p}_{i(t)} \rangle$ to A_j . (To take full advantage of locality, we must also set the weights d_{ji} so that players do not receive payments from adversaries they would otherwise not need to talk to.) So, by using targeted multicasts instead of broadcasts, we can confine each player’s messages to a small area of the network; in this case, no single node or link will see even $O(k + n)$ traffic. We can sometimes reduce bandwidth even further by combining messages as they flow through the network: for example, two resource consumption messages destined for A_j may be combined by adding their reported consumption values.

Finally, any implementation needs to make sure that the agents cannot gain by circumventing the mechanism: e.g., no player should find out another’s plan before committing to her own.⁵

⁴Technically, the agents could multicast rather than broadcast, so that, e.g., one player would never see another player’s messages, but in practice one would not expect this optimization to save much.

⁵One of the undesirable effects resulting if we would permit agents to wait until all other agents have sent their

4 Design goals

In designing our mechanism, we hope to ensure that individual, incentive-driven behavior leads to desirable system-wide properties. Here, we establish some useful guarantees for the *negotiation version* of our mechanism. The guarantees are **convergence to Nash equilibrium**, **budget balance**, **individual rationality**, and **efficiency**.

These guarantees follow from the fact that the learned negotiation outcome approaches a set of saddle-points of $\Omega(\mathbf{p}, \mathbf{a})$ in the limit (Thm. 3.4). By continuity of Ω , we can therefore conclude that, if we allow sufficient time for negotiation, the negotiation outcome is approximately a saddle-point. (We will not address distributed detection of convergence, but merely assume that we use our global regret bounds to calculate a sufficiently large T ahead of time; obviously efficiency could be improved by allowing early stopping.)

Convergence to Nash equilibrium. When working with selfish, strategic agents, we want to know whether a selfish agent has an incentive to unilaterally deviate from its part of the negotiation outcome. The following theorems show that the answer is, at least approximately, *no*: in the limit of large T , the negotiation outcomes $\bar{\mathbf{p}}_{[T]}, \bar{\mathbf{a}}_{[T]}$ converge to a subset of Nash equilibria. So, by continuity, $(\bar{\mathbf{p}}_{[T]}, \bar{\mathbf{a}}_{[T]})$ is an approximate Nash equilibrium for sufficiently large T —that is, each agent has a vanishing incentive to deviate unilaterally.

Theorem 4.1. *Let F_{P_i} denote the feasible set of player agent P_i ($i \in \{1, \dots, k\}$) and F_{A_j} denote the feasible set of adversarial agent A_j ($j \in \{1, \dots, n\}$). We have:*

$$\forall i \forall \mathbf{p}'_i \in F_{P_i} \forall \tilde{\mathbf{a}} \in K_A, \tilde{\mathbf{p}} \in K_P : \Omega_{P_i}(\mathbf{p}'_i, \tilde{\mathbf{a}}) \geq \Omega_{P_i}(\tilde{\mathbf{p}}_i, \tilde{\mathbf{a}}).$$

Proof. Let $\tilde{\mathbf{p}} \in K_P, \tilde{\mathbf{a}} \in K_A$. We know $(\tilde{\mathbf{p}}, \tilde{\mathbf{a}})$ is a saddle-point of Ω with respect to minimization over F_P and maximization over F_A . Hence, $\forall \mathbf{p}' \in F_P : \Omega(\tilde{\mathbf{p}}, \tilde{\mathbf{a}}) \leq \Omega(\mathbf{p}', \tilde{\mathbf{a}})$. Since $\Omega(\mathbf{p}, \mathbf{a}) = \sum_m \Omega_{P_m}(\mathbf{p}_m, \mathbf{a})$, $\forall \mathbf{p}, \mathbf{a}$, we have in particular: $\forall \mathbf{p}'_i \in F_{P_i} : \Omega_{P_i}(\mathbf{p}'_i, \tilde{\mathbf{a}}) + \sum_{m \neq i} \Omega_{P_m}(\tilde{\mathbf{p}}_m, \tilde{\mathbf{a}}) = \Omega(\mathbf{p}'_i, \tilde{\mathbf{p}}_{-i}, \tilde{\mathbf{a}}) \geq \Omega(\tilde{\mathbf{p}}, \tilde{\mathbf{a}}) = \Omega_{P_i}(\tilde{\mathbf{p}}_i, \tilde{\mathbf{a}}) + \sum_{m \neq i} \Omega_{P_m}(\tilde{\mathbf{p}}_m, \tilde{\mathbf{a}})$. Thus, $\Omega_{P_i}(\mathbf{p}'_i, \tilde{\mathbf{a}}) \geq \Omega_{P_i}(\tilde{\mathbf{p}}_i, \tilde{\mathbf{a}})$, $\forall \mathbf{p}'_i \in F_{P_i}$. $\square$

Theorem 4.2. *Let $(\tilde{\mathbf{p}}, \tilde{\mathbf{a}})$ be a saddle-point of $\Omega(\mathbf{p}, \mathbf{a})$ with respect to minimizing over $\mathbf{p}$ and maximizing over $\mathbf{a}$. We have: $\Omega_{A_j}(\tilde{\mathbf{a}}^j, \tilde{\mathbf{p}}) = \max_{a^j \in F_{A_j}} \Omega_{A_j}(a^j, \tilde{\mathbf{p}})$, $\forall j \in \{1, \dots, n\}$.*

Proof. Since $\sum_{j=1}^n \Omega_{A_j}(\tilde{\mathbf{a}}^j, \tilde{\mathbf{p}}) = \Omega_A(\tilde{\mathbf{a}}, \tilde{\mathbf{p}}) = \max_{\mathbf{a}} \Omega_A(\mathbf{a}, \tilde{\mathbf{p}}) = \max_{\mathbf{a} \in F_A} \sum_{j=1}^n \Omega_{A_j}(\tilde{\mathbf{a}}^j, \tilde{\mathbf{p}}) = \sum_{j=1}^n \max_{a^j \in F_{A_j}} \Omega_{A_j}(a^j, \tilde{\mathbf{p}})$, we have $\sum_{j=1}^n (\max_{a^j \in F_{A_j}} [\Omega_{A_j}(a^j, \tilde{\mathbf{p}})] - \Omega_{A_j}(\tilde{\mathbf{a}}^j, \tilde{\mathbf{p}})) = 0$. On the other hand, $\forall j : \max_{a^j \in F_{A_j}} [\Omega_{A_j}(a^j, \tilde{\mathbf{p}})] - \Omega_{A_j}(\tilde{\mathbf{a}}^j, \tilde{\mathbf{p}}) \geq 0$. Hence, $\forall j : \max_{a^j \in F_{A_j}} [\Omega_{A_j}(a^j, \tilde{\mathbf{p}})] - \Omega_{A_j}(\tilde{\mathbf{a}}^j, \tilde{\mathbf{p}}) = 0$. $\square$

Theorem 4.2 allows us to conclude that the individual part of an adversarial agent's negotiation outcome is a best-response action:

plans before choosing an own plan is that we could run into a deadlock. In general, we could enforce a commitment to a certain plan prior to learning about the other agents' current choices by employing cryptographic methods such as commitment schemes [Blu81, Eve82, Nao91].

Corollary 4.3. Let F_{A_j} denote the feasible set of adversarial agent A_j ($j \in \{1, \dots, n\}$) and F_{P_i} denote the feasible set of player agent P_i ($i \in \{1, \dots, k\}$). We have:

$$\forall j \forall a^j \in F_{A_j} \forall \tilde{\mathbf{p}} \in K_P, \tilde{\mathbf{a}} \in K_A : \Omega_{A_j}(a^j, \tilde{\mathbf{p}}) \leq \Omega_{A_j}(\tilde{a}^j, \tilde{\mathbf{p}}).$$

Budget balance. Since our overall goal is a socially optimal plan, we would hope that our mechanism neither siphons off money from the agents by running a surplus, nor requires continuous investment to fund a deficit. This is the question of budget balance. Since the agents make payments only to one another (and not directly to the mechanism), in one sense our mechanism is trivially budget balanced. However, a more interesting question is whether the mechanism is budget balanced if we consider the adversarial agents to be part of the mechanism—this additional property guarantees that the adversarial agents do not, in the long run, siphon off money or require external funding. Since we showed (in Sec. 3.1) that the adversarial agents each have zero revenue at any saddle point, and since the outcome of negotiation is an approximate saddle point, our mechanism is (approximately) budget balanced in this sense as well.

Budget balance can be evaluated *ex ante*, *ex interim*, or *ex post*, depending on whether it holds (in expectation) before the agents know their private information, after they know their private information but before they know the outcome of the mechanism, or after they know the outcome of the mechanism. Ex-post budget balance is the strongest property; our argument in fact shows approximate ex-post budget balance.

Individual rationality. Strategic agents will avoid participating in a mechanism if doing so improves their payoffs. A mechanism is individually rational if each agent is no worse off when joining the mechanism than when avoiding it. Just as with budget balance, we can speak of *ex-ante*, *ex-interim*, or *ex-post* individual rationality.

To make the question of individual rationality well-defined, we need to specify what happens if an agent avoids the mechanism. If an adversarial agent refuses to participate, we will assume that her corresponding resource goes unmanaged: no price is announced for it, and the player agents pay their natural costs for it. The adversarial agent therefore gets no revenue, either positive or negative. If a player agent refuses to participate, we will assume that she is constrained use no resources, that is, $\langle \mathbf{l}_{ji}, \mathbf{p}_i \rangle = 0$ for all j . (So, we assume that there is a plan satisfying these constraints.)

Since we showed that $\sup_{a^j} [\Omega_{A_j}(\mathbf{p}, a^j) - \beta_j(\nu_j(\mathbf{p}))] = 0$ (in Sec. 3.1), A_j has (approximately) no incentive to avoid the mechanism when we play an (approximate) saddle point. So, the mechanism is approximately ex-post IR for adversaries.

If a player agent does not participate in the mechanism, she has no chance of acquiring any resources. Since she would not have to pay for joining and using no resources (the remainder $r_j(a^j)$ is nonnegative), it is irrational not to join. So, the mechanism is ex-post IR for players.

Efficiency. A mechanism is called *efficient* if its outcome minimizes global social cost. Thm. 3.4 showed that the mechanism finds an approximate saddle-point of Ω . We showed in Sec. 3.1 that, in any saddle-point, the player cost is $\min_{\mathbf{p}} \omega(\mathbf{p})$ (the socially-optimal cost), and the adversary cost is 0. So, in an approximate saddle-point, the social cost is approximately optimal; the mechanism is therefore approximately efficient.

5 Related Work

The idea of using no-regret algorithms to solve OCPs in order to accomplish a planning task is not new (e.g., [BBCM03]). It has, for instance, been proposed for online routing in the Wardrop setting of multi-commodity flows [BEDL06], where the authors established convergence to Nash equilibrium for infinitesimal agents. In contrast with this line of work, we seek *globally* good outcomes, rather than just equilibria.

Another body of related work is concerned with selfish routing in *nonatomic* settings (with infinitesimal agents—e.g., [BEDL06, Rou07]). Many of these works provide strong performance guarantees and price of anarchy results considering selfish agents. We consider a similar but not identical setup, with a finite number of agents and divisible resources.

As mentioned before, our planning approach can be given a simple market interpretation: interaction among player agents happens indirectly through resource prices learned by the adversaries. Many researchers have demonstrated experimental success for market-based planners (e.g., [SD99, GM02, SDZK04, Wel93, GK07, MWY04]). While these works experimentally validate the usefulness of their approaches and implement distributivity, only a few provide guarantees of optimality or approximate optimality (e.g., [LMK⁺05, Wel93]).

Guestrin and Gordon proposed a decentralized planning method using a distributed optimization procedure based on Benders decomposition [GG02]. They showed that their method would produce approximately optimal solutions and offered bounds to quantify the quality of this approximation. However, as with most authors, they assumed agents to be *obedient*, i.e., to follow the protocol in every aspect. By contrast, we address *strategic* agents, i.e., selfish, incentive-driven entities prone to deviating from prescribed behavior if it serves their own benefit. But, since the trick of dualizing constraints to decouple an optimization problem is analogous to Benders decomposition, we can view our mechanism as a generalization of Guestrin and Gordon’s method to decentralized computation on selfish agents.

Designing systems that provably achieve a desired global behavior with strategic agents is exactly the field of study of classic mechanism design. Many mechanisms, though, are heavy-weight and centralized, and are concerned neither with distributed implementation nor with computational feasibility. Attempting to fill this gap, a new strand of work under the label *distributed algorithmic mechanism design* has evolved [FPS01, FPSS05, Wel93].

Our approach combines many advantages of the above branches of work for multiagent planning. It is distributed, and provides asymptotic guarantees regarding mechanism design goals such as budget balance and quality of the learned solution. If we consider the adversarial agents to be independent, selfish entities that are not part of the mechanism, the proposed mechanism is relatively light-weight; it merely offers infrastructure for the participating agents to coordinate their planning efforts through learning in a repeated game. And, as the following section shows, its theoretical guarantees translate into reliable practical performance, at least in our small-scale network routing experiments.

6 Experiments

We conducted experiments on a small multi-agent min-cost routing domain. We model our network by a finite, directed graph with edges E (physical links) and vertices V (routers). Each edge $e \in E$ has a finite capacity $\gamma(e)$, as well as a fixed intrinsic cost c^e for each unit of traffic routed through e . We assume that bandwidth is infinitely divisible.

Players are indexed by a source vertex s and a destination vertex r . Player P_{sr} wants to send an amount of flow d_{sr} from s to r . P_{sr} 's individual plan is a vector $\mathbf{f}_{sr} = (f_{sr}^e)_{e \in E}$, where $f_{sr}^e \in [0, U]$ is the amount of traffic that P_{sr} routes through edge e . (U is an upper bound, chosen ahead of time to be larger than the largest expected flow.) Feasible plans are those that satisfy *flow conservation*, i.e., the incoming traffic to each vertex must balance the outgoing traffic. We also excluded plans which route flow in circles.

If the total usage of an edge e exceeds its capacity, all agents experience an increased cost for using e . This extra cost could correspond to delays, or to surcharges from a network provider. We set the global penalty for edge e to be a hinge loss function $\beta_e(\nu) = \max\{0, u_e \nu\}$, so the total cost of e increases linearly with the amount of overuse. For convenience we also modeled each player's demand d_{sr} as a soft constraint $\beta_{sr}(\nu) = \max\{0, u_{sr} \nu\}$, although doing so is not necessary to achieve a distributed mechanism.

Applying our problem transformation led to new, total player cost $\Omega(\mathbf{f}, \mathbf{a}) = \sum_{s,r} \sum_{e \in E} c^e f_{sr}^e + \sum_{e \in E} \Omega_{A_e}(a_{\text{cap}}^e, \mathbf{f}) + \sum_{s,r} \Omega_{A_{sr}}(a_d^{sr}, \mathbf{f})$ in the mechanism. Here, for each edge e , we introduced an adversarial agent A_e , who controls the cost for capacity violations at e by setting the price a_{cap}^e . And, as a slight extension to our general description in Section 3, we introduced additional adversarial agents to implement the soft constraints on demand; agent A_{sr} chooses a price a_d^{sr} for failing to meet demand on route sr .

With this setup, we ran more than 2800 simulations, for the most part on random problem instances, but also for manually-designed problems on graphs of sizes varying between 2 and 16 nodes. In each instance there were between 1 and 32 player agents. For no-regret learning, we used the Greedy Projection algorithm [Zin03] with $(\frac{1}{\sqrt{t}})_{t \in \mathbb{N}}$ as the sequence of learning rates.

A simple example of an averaged player plan after a number of iterations is depicted in Figure 2(b). In this experiment, we had a 6-node network and three players $P_{2,3}$, $P_{1,4}$, and $P_{4,6}$ with demands 30, 70 and 110. We set $c^{(2,3)}, c^{(3,2)} = 10$, and $c^e = 1$ for all other edges e . Edges $(5, 6)$ and $(6, 5)$ had capacities of 50, while all other capacities were 100. Our method successfully discovered that $P_{4,6}$ should send as much flow as possible through the cheap edge $(5, 6)$, and the rest along the expensive path through $(3, 2)$. Adversarial agents successfully discouraged the players from violating the capacity constraints, while simultaneously making sure that as much demand as possible was satisfied. Also note that player $P_{2,3}$ served the common good by (on average) routing flow through the pricey edge $(2, 3)$ instead of taking the path through the bottleneck $(6, 5)$; this latter path would have been cheaper for $P_{2,3}$ if we didn't consider the extra costs imposed by the adversarial agents. The plots in Figs. 1 and 2 validate our theoretical results: Figs. 1(a),(b) demonstrate that the regrets of the combined agents P and A converge to zero, as shown in Lem. 3.2 and Lem. 3.3. Fig. 2(a) demonstrates convergence to a saddle-point of Ω . In the plot, the upper curve shows $\max_{\mathbf{a}} \Omega(\bar{\mathbf{f}}_{[T]}, \mathbf{a})$, while the lower curve shows $\min_{\mathbf{f}} \Omega(\mathbf{f}, \bar{\mathbf{a}}_{[T]})$. The horizontal, dashed line is the minimax value of Ω (which was 30 in the corresponding experiment). As guaranteed by Thm. 3.4,

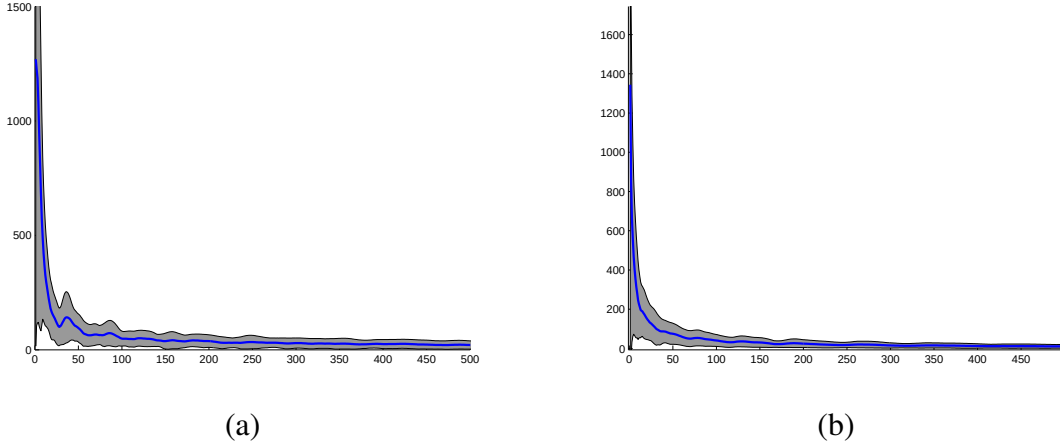


Figure 1: Average positive regret of synthesized player (a) and synthesized adversary (b), averaged over 100 random problems, as a function of iteration number. Grey area indicates standard error.

the three curves converge to one another. While the fact of convergence in these figures is not a surprise, it is reassuring to see that the convergence is fast in practice as well as in theory.

7 Discussion

We presented a distributed learning mechanism for use in multiagent planning. The mechanism works by introducing *adversarial* agents who set taxes on common resources. By so doing, it *decouples* the original player agents’ planning problems. We then proposed that the original and adversarial agents should learn about one another by playing the decoupled planning game repeatedly, either in reality (the *online* setup) or in simulation (the *negotiation* setup).

We established that, if all agents use no-regret learning algorithms in this repeated game, several desirable properties result. These properties included convergence of $\bar{\mathbf{p}}_{[T]}$, the average composite plan, to a socially optimal solution of the original planning problem, as well as convergence of $\bar{\mathbf{p}}_{[T]}$ and the corresponding adversarial tax-plan $\bar{\mathbf{a}}_{[T]}$ to a Nash equilibrium of the game. We also showed that our mechanism is budget-balanced in the limit of large T .

So far, we do not know in what cases our mechanism is incentive-compatible; in particular, we do not know when it is rational for the individual agents to employ no-regret learning algorithms. Certainly, we can invent cases where it is not rational to choose a no-regret algorithm, but we believe that there are practical situations where no-regret algorithms are a good choice. Investigating this matter, and modifying the mechanism to ensure incentive compatibility in all cases, is left to future work.

Compared to a centralized planner, our method can greatly reduce the bandwidth needed at the choke-point agent. (The choke-point agent is the one who needs the most bandwidth; in a centralized approach it is normally the centralized planner.) In very large systems, agents P_i and A_j only need to send messages to one another if P_i considers using resource j , so we can often use locality constraints to limit the number of messages we need to send.

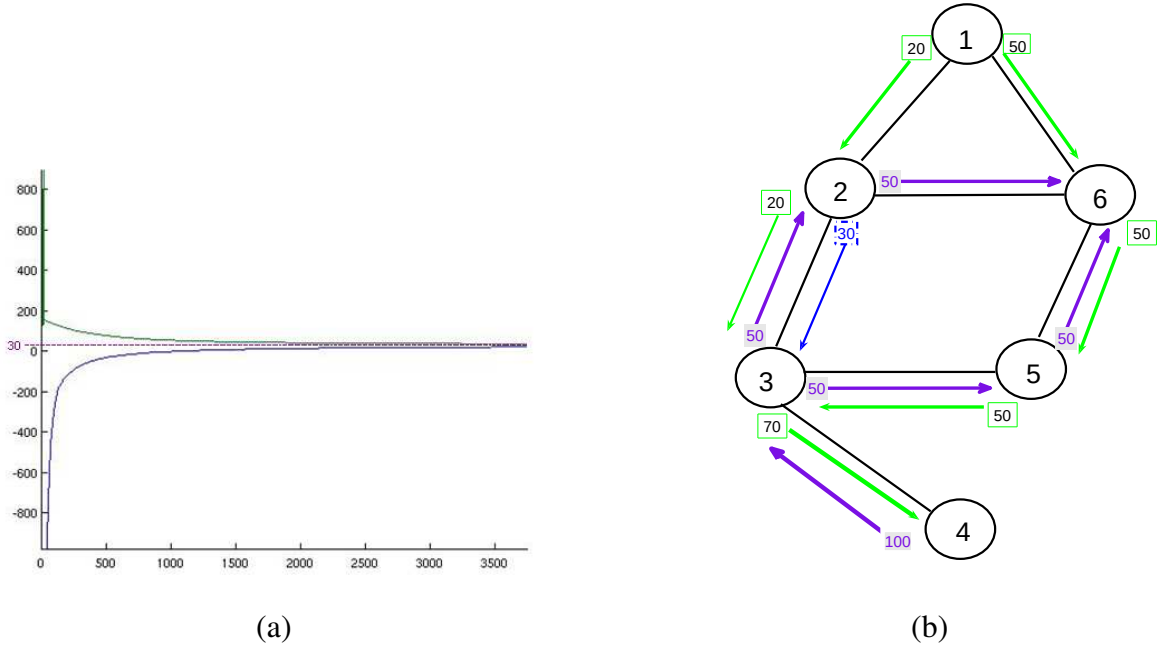


Figure 2: (a): Payoff Ω for the synthesized player (upper curve) and adversary (lower curve) when P and A play their averaged strategy against a best-response opponent in each iteration. Horizontal line shows minimax value. (b): Average plan for three agents in a 6-node instance after 10000 iterations, rounded to integer flows. The displayed plan is socially optimal.

Our method combines desirable features from various previous approaches: like centralized mechanisms and some other distributed mechanisms we can provide rigorous guarantees such as social optimality and individual rationality. But, like prior work in market-based planning, we expect our approach to be efficient and implementable in a distributed setting. Our experiments tend to confirm this prediction.

Acknowledgements

We would like to thank Ram Ravichandran for his kind assistance with proofreading this technical report. This research was funded in part by a grant from DARPA’s Computer Science Study Panel program. All opinions and conclusions are the authors’.

References

- [BBCM03] N. Bansal, A. Blum, S. Chawla, and A. Meyerson. Online oblivious routing. In *SPAA '03: Proceedings of the fifteenth annual ACM symposium on Parallel algorithms and architectures*, pages 44–49, 2003.
- [BEDL06] Avrim Blum, Eyal Even-Dar, and Katrina Ligett. Routing without regret: on convergence to nash equilibria of regret-minimizing algorithms in routing games. In *PODC '06: Proceedings of the twenty-fifth annual ACM symposium on Principles of distributed computing*, pages 45–52, 2006.
- [BHLR07] Avrim Blum, Mohammad Taghi Hajiaghayi, Katrina Ligett, and Aaron Roth. Regret minimization and the price of total anarchy. Working paper, 2007.
- [Blu81] Manuel Blum. Coin flipping by telephone. In *Crypto '81*, 1981.
- [BV04] Stephen Boyd and Lieven Vandenberghe. *Convex Optimization*. Cambridge University Press, 2004.
- [CG08] Jan-P. Calliess and Geoffrey J. Gordon. No-regret learning and a mechanism for distributed multiagent planning. In *Proc. of 7th Int. Conf. on Autonomous Agents and Multiagent Systems (AAMAS'08)*, 2008.
- [Eve82] S. Even. A protocol for signing contracts. Technical Report Tech. Rep. 231, Technion, Haifa, Israel, 1982.
- [FPS01] Joan Feigenbaum, Christos H. Papadimitriou, and Scott Shenker. Sharing the cost of multicast transmissions. *J. Comput. Syst. Sci.*, 63(1):21–41, 2001.
- [FPSS05] Joan Feigenbaum, Christos Papadimitriou, Rahul Sami, and Scott Shenker. A bgp-based mechanism for lowest-cost routing. *Distrib. Comput.*, 18(1):61–72, 2005.
- [FS96] Yoav Freund and Robert E. Shapire. Game theory, on-line prediction and boosting. In *COLT*, 1996.
- [GG02] C. Guestrin and G. Gordon. Distributed planning in hierarchical factored mdps. In *UAI*, 2002.
- [GGMZ07] Geoffrey J. Gordon, Amy Greenwald, Casey Marks, and Martin Zinkevich. No-regret learning in convex games. Technical Report CS-07-10, Brown University, 2007.
- [GK07] E. Gomes and R. Kowalczyk. Reinforcement learning with utility-aware agents for market-based resource allocation. In *AAMAS*, 2007.
- [GM02] B. Gerkey and M. Mataric. Sold!: Auction methods for multirobot coordination. *IEEE Transactions on Robotics and Automation*, 19(5):758–768, 2002.

- [Gor99] Geoffrey J. Gordon. Regret bounds for prediction problems. In *COLT: Workshop on Computational Learning Theory*, 1999.
- [LMK⁺05] M. Lagoudakis, V. Markakis, D. Kempe, P. Keskinocak, S. Koenig, A. Kleywegt, C. Tovey, A. Meyerson, and S. Jain. Auction-based multi-robot routing. In *Proceedings of the International Conference on Robotics: Science and Systems (ROBOTICS)*, pages 343–350, 2005.
- [MJHA05] Gerhard Illing Manfred J. Holler (Author). *Einführung in die Spieltheorie*. Springer Akad. Verlag, 2005.
- [MWY04] N. Muguda, P. R. Wurman, and R. M. Young. Experiments with planning and markets in multi-agent systems. *SIGecom Exchanges*, 5:34–47, 2004.
- [Nao91] M. Naor. Bit commitment using pseudorandomness. *Journal of Cryptology: the journal of the International Association for Cryptologic Research*, 4(2):151–158, 1991.
- [Owe95] G. Owen. *Game Theory*. Academic Press, 1995.
- [Roc70] R. Tyrell Rockafellar. *Convex Analysis*. Princeton University Press, 1970.
- [Rou07] T. Roughgarden. Selfish routing and the price of anarchy (survey). In *OPTIMA*, 2007.
- [SD99] A. Stentz and M. B. Dias. A free market architecture for coordinating multiple robots. Technical Report CMU-RI-TR-99-42, Carnegie Mellon, Robotics Institute, Pittsburgh, PA, USA, 1999.
- [SDZK04] Anthony (Tony) Stentz, M. Bernardine Dias, Robert Michael Zlot, and Nidhi Kalra. Market-based approaches for coordination of multi-robot teams at different granularities of interaction. In *Proc. ANS 10th Int. Conf. on Robotics and Rem. Sys. for Hazardous Env.*, 2004.
- [SL07] G. Stoltz and G. Lugosi. Learning correlated equilibria in games with compact sets of strategies. *Games and Economic Behavior*, 59:187–208, 2007.
- [SLBon] Y. Shoham and K. Leyton-Brown. *Multiagent Systems- Draft of Dec. 2006*. Cam. Univ. Press, In preparation.
- [vN28] J. von Neumann. Zur theorie der gesellschaftsspiele. *Mathematische Annalen*, 100:295–320, 1928.
- [Wel93] M. P. Wellman. A market-orient programming environment and its application to distributed multi-commodity flow problems. *J. of Artificial Intelligence Research*, 1:1–23, 1993.
- [Zin03] M. Zinkevich. Online convex programming and generalized infinitesimal gradient ascent. In *Twentieth International Conference on Machine Learning*, 2003.

A Background

A.1 Convex Sets and Convex Functions

In this subsection we will list a couple of basic definitions and theorems regarding convex functions for later use. The material was extracted from [BV04] and [Roc70].

Definition A.1 (Affine Function). A mapping ϕ between vector spaces V_1 and V_2 is called *affine* iff $\exists \alpha \in \text{hom}(V_1, V_2), \mathbf{r} \in V_2 \forall \mathbf{v}_1 \in V_1 : \phi(\mathbf{v}_1) = \alpha(\mathbf{v}_1) + \mathbf{r}$.

Definition A.2 (Affine Sets). $A \subset \mathbb{R}^d$ is $:\Leftrightarrow \forall \mathbf{a}_1, \mathbf{a}_2 \in A \forall l \in \mathbb{R} : (1-l)\mathbf{a}_1 + l\mathbf{a}_2 \in A$.

We can see that an affine set is a line through any two of its points by rewriting $(1-l)\mathbf{a}_1 + l\mathbf{a}_2$ as $\mathbf{a}_1 + l(\mathbf{a}_2 - \mathbf{a}_1)$. This idea can be generalized to multiple points leading to the idea of an *affine hull*.

Definition A.3 (Affine Hull). Let $S \subset \mathbb{R}^d$ be a subset of $\mathbb{R}^d$. Its *affine hull* $\text{aff } S$ is the smallest (with respect to the inclusion operator $\subset$) affine set containing S .

Note, $\text{aff } S = \{\sum_{i=1}^n r_i \mathbf{s}_i | \mathbf{s}_i \in S, \sum_{i=1}^n r_i = 1, n \in \mathbb{N}\}$ (cf. [Roc70]).

Instead of asking that for any two points of a set the whole line through them is contained in the set, we could be interested in sets with the property that only all points on the line segment between any two of its points need to be contained as well. This leads to the important notion of convex sets.

Definition A.4 (Convex Sets). A set of vectors $C \subset \mathbb{R}^d$ is convex $:\Leftrightarrow \forall \mathbf{c}_1, \mathbf{c}_2 \in C \forall l \in [0, 1] : (1-l)\mathbf{c}_1 + l\mathbf{c}_2 \in C$.

As an example, we will describe the *polyhedrons*. Polyhedrons are solution sets of a collection of linear inequalities and equalities [BV04]. An example for a polyhedron is the set $\{\mathbf{v} \in \mathbb{R}^d | \langle \mathbf{l}_1, \mathbf{v} \rangle \leq y_1, \dots, \langle \mathbf{l}_n, \mathbf{v} \rangle \leq y_n, \langle \mathbf{e}, \mathbf{v} \rangle = x\}$. Bounded polyhedrons are usually called *polytopes*. The convex hull of a set (the intersection of all convex supersets) is a polyhedron [Roc70]. Following standard literature we will usually consider functions ϕ of the form $\phi : S \subset \mathbb{R}^d \rightarrow \mathbb{R} \cup \{-\infty, +\infty\}$. We follow the customary convention to understand by $-\infty, +\infty$ mathematical objects with the property $\forall x \in \mathbb{R} : -\infty < x < +\infty$. By this point of view a distinction between ϕ 's *domain* S and its *effective domain* $\text{dom } \phi := \{\mathbf{s} \in S | \phi(\mathbf{s}) < +\infty\}$ is obtained. Instead of considering functions on S we will extend them to all of $\mathbb{R}^d$ by setting their values at all points in $\mathbb{R}^d$ not in S to $+\infty$. This has technical benefits which will not become apparent in this work but it is done to make our assumptions consistent with the ones needed for theorems we will borrow from [Roc70].

As we will see, another way to define the effective domain of a function is by its *epigraph* – the set of all points lying on or above its graph.

Definition A.5 (Epigraph). The epigraph $\text{epi } \phi$ of a function $\phi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{-\infty, +\infty\}$ is $\text{epi } \phi := \{(\mathbf{s}, r) \in \mathbb{R}^d \times \mathbb{R} | r \geq \phi(\mathbf{s})\}$.

Note, points in the epigraph only have real components, i.e. points with $-\infty$ or $+\infty$ components are not included. We can now define $\text{dom } \phi := \{\mathbf{s} \in \mathbb{R}^d \mid \exists r \in \mathbb{R} : (\mathbf{s}, r) \in \text{epi } \phi\}$.

Definition A.6 (Convex Function). A real-valued function is convex iff its epigraph is a convex set.

Definition A.7 (Concave Function). A real-valued function is called concave iff $-f$ is convex.

Theorem A.8 (Jensen's Inequality for Convex Functions (Cf. [BV04])). *If ϕ is a convex function and $z_1, \dots, z_n$ are positive real numbers such that $z_1 + \dots + z_n = 1$, then:*

$$\forall \mathbf{x}_1, \dots, \mathbf{x}_n \in \text{dom } \phi : \phi(z_1 \mathbf{x}_1 + \dots + z_n \mathbf{x}_n) \leq z_1 \phi(\mathbf{x}_1) + \dots + z_n \phi(\mathbf{x}_n).$$

By multiplying both sides of the inequality by -1 we can conclude:

Corollary A.9 (Jensen's Inequality for Concave Functions). *If ϕ is a concave function and $z_1, \dots, z_n$ are positive real numbers such that $z_1 + \dots + z_n = 1$, then:*

$$\forall \mathbf{x}_1, \dots, \mathbf{x}_n \in \text{dom } \phi : \phi(z_1 \mathbf{x}_1 + \dots + z_n \mathbf{x}_n) \geq z_1 \phi(\mathbf{x}_1) + \dots + z_n \phi(\mathbf{x}_n).$$

Remark A.10. It is to be noted that an affine function is both convex and concave. In particular, Jensen's inequality holds in both directions, i.e. Jensen's inequality is actually an equality for affine functions.

Remark A.11. $\text{dom } \phi$ results from the epigraph by projection on $\mathbb{R}^d$ which is a linear transformation. Since it is known that linear transformations conserve convexity (cf. [Roc70]) the effective domain of a convex function must be convex.

Theorem A.12 (Cf. [Roc70]). *The pointwise supremum of an arbitrary collection of convex functions is convex.*

Proof. The proof is following [Roc70]. Let I denote an index set, $\phi_i (i \in I)$ a collection of convex functions. $\phi(\mathbf{x}) := \sup \{\phi_i(\mathbf{x}) \mid i \in I\} \Rightarrow \text{epi } \phi = \bigcap_i \text{epi } \phi_i$. It is known that the intersection of convex sets results in a convex set again. $\square$

In order to be able to state the next theorem we will need to introduce more terminology.

Definition A.13 (Proper Convex Function). A convex function ϕ is called *proper* iff $\text{dom } \phi \neq \emptyset$ and $\forall \mathbf{x} \in \text{dom } \phi : \phi(\mathbf{x}) > -\infty$.

We could say proper convex functions represent the non-pathological case, they are the type of functions we would normally consider.

Definition A.14 (Relative Interior). Let C be a convex set. Its relative interior $\text{ri } C$ is defined as its interior points relative to $\text{aff } C$: $\text{ri } C = \{\mathbf{x} \in C \mid \exists r > 0 : B(\mathbf{x}, r) \cap \text{aff } C \subset C\}$.

The motivation for introducing the concept of relative interior points can be best understood by an example. Consider a line segment in $\mathbb{R}^3$. It does not have truly interior points in the whole metric space $\mathbb{R}^3$. However the points between its two delimiting end points are interior points relative to its affine hull $\mathbb{R}^1$.

Note, if we denote the closure of C by $\text{cl } C$ we have : $\text{ri } C \subset C \subset \text{cl } C \subset \text{cl } (\text{aff } C) = \text{aff } C$ ([Roc70]).

Theorem A.15. Let ϕ be a proper convex function and let S be any closed, bounded subset of $\text{ri}(\text{dom } \phi)$. Then ϕ is Lipschitzian relative to S .

Proof. Confer to [Roc70]. □

A.2 Saddle Points and Game-Theoretic Essentials

This work uses some game-theoretic notions that will be briefly listed in this section. For a proper game-theoretic introduction the reader is advised to refer to related textbooks such as [Owe95, MJHA05]. The game-theoretic terminology we need for our exposition is very basic and furthermore, our reduction to the synthesized agents allows us to restrict our attention to the case of two-player games. A central aspect in two-player games are minimax-equilibria which correspond to the general notion of saddle-points.

A.2.1 Saddle-points

Minimax theory treats a class of optimization problems which involve not just maximization or minimization, but a combination of the two. We will confine our considerations to the case of continuous functions on compact domains. A treatment of the more general case can be found in [Roc70]. Let $\phi : X \times Y \rightarrow \mathbb{R}$ where X, Y are compact sets.

We can consider a minimization problem for function $\max_{y \in Y} \phi(\mathbf{x}, \cdot) : X \rightarrow \mathbb{R}$ and a maximization problem for function $\min_{x \in X} \phi(\cdot, y) : Y \rightarrow \mathbb{R}$.

It is well known, that we always have $\max_{y \in Y} \min_{x \in X} \phi(\mathbf{x}, \mathbf{y}) \leq \min_{x \in X} \max_{y \in Y} \phi(\mathbf{x}, \mathbf{y})$.

If the values $\min_{x \in X} \max_{y \in Y} \phi(\mathbf{x}, \mathbf{y})$ and $\max_{y \in Y} \min_{x \in X} \phi(\mathbf{x}, \mathbf{y})$ coincide, the common value v is called *minimax- or saddle-value* of function ϕ (with respect to minimizing over X and maximizing over Y).

Definition A.16 (Saddle-point). A point $(\tilde{\mathbf{x}}, \tilde{\mathbf{y}}) \in X \times Y$ is a *saddle-point* (with respect to minimizing over X and maximizing over Y) if

$$\forall \mathbf{x} \in X \forall \mathbf{y} \in Y : \phi(\tilde{\mathbf{x}}, \mathbf{y}) \leq \phi(\tilde{\mathbf{x}}, \tilde{\mathbf{y}}) \leq \phi(\mathbf{x}, \tilde{\mathbf{y}}).$$

An adaptation of Lemma 36.2 in [Roc70] is:

Lemma A.17. $(\tilde{\mathbf{x}}, \tilde{\mathbf{y}}) \in X \times Y$ is a saddle-point of ϕ (with respect to minimizing over X and maximizing over Y) iff $\tilde{\mathbf{x}} = \arg \min_{x \in X} \max_{y \in Y} \phi(\mathbf{x}, \mathbf{y})$, $\tilde{\mathbf{y}} = \arg \max_{y \in Y} \min_{x \in X} \phi(\mathbf{x}, \mathbf{y})$ and the minimax value v exists. If $(\tilde{\mathbf{x}}, \tilde{\mathbf{y}})$ is a saddle-point then $\phi(\tilde{\mathbf{x}}, \tilde{\mathbf{y}}) = v$.

A.2.2 Some fundamental game-theoretic notions

Many game-theoretic problems can be reduced to a specific type of game called *normal – form* game.

Definition A.18 (cf. [SLBon]). A finite n -player normal form game is a tuple $(P, N, A, O, \mu, \varpi)$, where P is a finite set of k players indexed by i , $A = (A_1, \dots, A_k)$, where A_i is a finite set of actions (or synonymously *pure strategies*) available to the i th player. Each $a = (a_1, \dots, a_n) \in A$

is called action profile. O denotes a set of outcomes, $\mu : A \rightarrow O$ maps an action to an outcome, $\varpi = (\varpi_1, \dots, \varpi_k)$ is the ordered set of individual payoff functions, where $\varpi_i : O \rightarrow \mathbb{R}$ determines the individual payoff for the i th player.

Note, a more wide-spread definition of normal-form games avoids the introduction of the outcome set O and directly maps actions to payoffs. We can consider this as a special case of our definition by setting $O := A$ and μ to the identity mapping. Unless otherwise stated, we assume to work with this latter configuration. Also note, we can easily extend the definition to handle infinite (but compact) action sets as will be required for our treatment.

Agents are generally not required to always play the same, fixed action. Instead they may be allowed to randomize over their action set. A distribution over an individual action set A_i is called *mixed strategy*. A pure strategy can be considered as a special case of a mixed strategy. Many times, agents playing actions generated according to mixed strategies can have higher expected payoffs than if they would restrict themselves to deterministic behavior. The ordered set of all mixed strategies $s = (s_1, \dots, s_k)$ of all agents is called *strategy profile*.

In game theory, it is standard to assume agents playing the game act rational, i.e. their behavior is governed by the desire to maximize their individual payoff functions. Given the other agents' strategy profile s_{-i} it is in the best interest of the i th rational player to play a *best response* s_i^* to it. The best response s_i^* is given by $s_i^* = \operatorname{argmax}_{s_i} \varpi_i(s_i, s_{-i})$.

Certain types strategy profiles that can happen to be played are of special interest to strategic agents. A famous one is the so-called *Nash-equilibrium*.

Definition A.19 (Nash-equilibrium). A strategy profile $\tilde{s} = (\tilde{s}_1, \dots, \tilde{s}_k)$ is a Nash-equilibrium, if for each $i \in \{1, \dots, k\}$ we have $\varpi_i(s_i, \tilde{s}_{-i}) \leq \varpi_i(\tilde{s}_i, \tilde{s}_{-i})$.

In other words, a Nash-equilibrium is a strategy profile where each player plays best-response to the strategies of all other players. Hence, if every player knew that the current strategy profile is a Nash-equilibrium, no single player would have an incentive to unilaterally deviate from its current strategy.

As a special case, we now consider two-player (normal-form) games. Such a game is called *zero-sum*, if $\varpi_1(s_1, s_2) = -\varpi_2(s_1, s_2) \forall s_1, s_2$. Therefore, we can redefine say player 1's desire to maximize its payoff by stating that its goal is to minimize its loss given by $\varpi := \varpi_2$ instead.

A famous result, called *Minimax-Theorem* or *von Neumann's Theorem*, found by J. von Neumann states that in every Nash-equilibrium of a finite, two-player zero-sum game the minmax-value is attained and equals the maxmin-value and that in every such a game, a Nash-equilibrium exists.

So, a Nash-equilibrium in a two-player zero-sum game is given by a saddle-point $(\tilde{s}_1, \tilde{s}_2)$ of ϖ (which directly follows from Definition A.16 or with the Minimax-Theorem in conjunction with Lemma A.17). Since $\varpi(\tilde{s}_1, \tilde{s}_2)$ equals the minimax value such a saddle-point forming the Nash-equilibrium is commonly referred to as a *minimax-equilibrium* of the game.

B Supplementary prerequisites

Next, we will establish some properties needed in Appendix C. We do not suppose they are novel since the material seems standard. However, we provide our own derivations.

Lemma B.1. *Let $\Omega : X \times Y \rightarrow \mathbb{R}$ be a continuous mapping (on $X \times Y$) where X, Y are **compact** subsets of Hilbert-spaces. We have:*

$\Psi : Y \rightarrow \mathbb{R}, \mathbf{y} \mapsto \inf_{\mathbf{x} \in X} \Omega(\mathbf{x}, \mathbf{y})$ is continuous.

Proof. Due to compactness we have $\Psi(\mathbf{y}) = \min_{\mathbf{x} \in X} \Omega(\mathbf{x}, \mathbf{y}), \forall \mathbf{y}$. Let $(\mathbf{y}_n)$ be a sequence in Y converging to $\mathbf{y} \in Y$ as $n \rightarrow \infty$. We wish to show that $\lim_{n \rightarrow \infty} \Psi(\mathbf{y}_n) = \Psi(\mathbf{y})$. Before we proceed with the proof we need to establish two prerequisites.

(i) Let $(\mathbf{a}_n)$ be a sequence in H and $\mathbf{a} \in H$ where H is a Hilbert-space with norm $\|\cdot\|$.

CLAIM1: If every subsequence of $(\mathbf{a}_n)$ contains a subsequence converging to $\mathbf{a}$ then $(\mathbf{a}_n)$ itself converges to $\mathbf{a}$.

Proof (CLAIM1): Let $\mathbf{a}_n \not\rightarrow \mathbf{a}$. Hence, $\exists e \geq 0 \forall n \in \mathbb{N} \exists m(n) \geq n : \|\mathbf{a}_{m(n)} - \mathbf{a}\| \geq e$. Then we can define the subsequence $(\mathbf{a}_{m(n)})$ where $\|\mathbf{a}_{m(n)} - \mathbf{a}\| \geq e, \forall n \in \mathbb{N}$. This subsequence of $(\mathbf{a}_n)$ does not contain a subsequence converging to $\mathbf{a}$. *q.e.d.*

(ii) Let $(\mathbf{y}_n)$ be a sequence in Y converging to $\mathbf{y} \in Y$ as $n \rightarrow \infty$ and (a_n) be the sequence in $\mathbb{R}$ defined as $a_n = \Psi(\mathbf{y}_n), \forall n \in \mathbb{N}$.

CLAIM2: There exists a subsequence of (a_n) converging to $\Psi(\mathbf{y})$.

Proof (CLAIM2): Define an arbitrary sequence $(\mathbf{x}_n)$ such that $\Omega(\mathbf{x}_n, \mathbf{y}_n) = \min_{\mathbf{x} \in X} \Omega(\mathbf{x}, \mathbf{y}_n), \forall n \in \mathbb{N}$. Note, such a sequence exists (owing to compactness of X) and we have $\forall n \in \mathbb{N} : \Psi(\mathbf{y}_n) = \Omega(\mathbf{x}_n, \mathbf{y}_n)$.

Since $X \times Y$ is a compact set in a complete metric space, there exists a convergent subsequence $(\mathbf{x}'_n, \mathbf{y}'_n)$ of sequence $(\mathbf{x}_n, \mathbf{y}_n)$. Let $\mathbf{x} \in X$ such that $(\mathbf{x}'_n, \mathbf{y}'_n) \rightarrow (\mathbf{x}, \mathbf{y})$.

We have : $\forall \mathbf{v} \in X : \Omega(\mathbf{v}, \mathbf{y}) = \lim_{n \rightarrow \infty} \Omega(\mathbf{v}, \mathbf{y}'_n) \geq^6 \lim_{n \rightarrow \infty} \Omega(\mathbf{x}'_n, \mathbf{y}'_n) =^7 \Omega(\mathbf{x}, \mathbf{y})$. Hence, $\Omega(\mathbf{x}, \mathbf{y}) = \min_{\mathbf{v} \in X} \Omega(\mathbf{v}, \mathbf{y}) = \Psi(\mathbf{y})$.

Define (a'_n) where $a'_n := \Psi(\mathbf{y}'_n)$. We know (a'_n) is a subsequence of (a_n) . Now, we have $\lim_{n \rightarrow \infty} a'_n = \lim_{n \rightarrow \infty} \Psi(\mathbf{y}'_n) = \lim_{n \rightarrow \infty} \Omega(\mathbf{x}'_n, \mathbf{y}'_n) = \Omega(\mathbf{x}, \mathbf{y}) = \Psi(\mathbf{y})$. *q.e.d.*

The final argument completing the proof is the following: If (a'_n) is any subsequence of $(a_n) = (\Psi(\mathbf{y}_n))$ then there is a subsequence $(\mathbf{y}'_n)$ of $(\mathbf{y}_n)$ still converging towards $\mathbf{y}$ such that $a'_n = \Psi(\mathbf{y}'_n), \forall n \in \mathbb{N}$. Sequence (a'_n) meets the preconditions of (ii) and thus, we know that there is a subsequence of subsequence (a'_n) which converges towards $a = \Psi(\mathbf{y})$.

Hence, every subsequence of (a_n) contains a subsequence that converges to a . By (i), we know that $\Psi(\mathbf{y}_n) = a_n \rightarrow a = \Psi(\mathbf{y})$ as $n \rightarrow \infty$. □

⁶Since by definition $\mathbf{x}'_n \in \operatorname{argmin}_{\mathbf{x} \in X} \Omega(\mathbf{x}, \mathbf{y}'_n)$.

⁷By continuity.

Providing an analogous argument one can prove the following lemma:

Lemma B.2. *Let $\Omega : X \times Y \rightarrow \mathbb{R}$ be a continuous mapping (on $X \times Y$) where X, Y are **compact** subsets of Hilbert-spaces. We have:*

$\phi : X \rightarrow \mathbb{R}, \mathbf{x} \mapsto \sup_{\mathbf{y} \in Y} \Omega(\mathbf{x}, \mathbf{y})$ is continuous.

Proof. The proof is completely analogous to the one provided for Lemma B.1 and shall be omitted here. $\square$

C Derivation of Theorems 3.4 and 3.5

In Section 3.2.2, we presented Theorems 3.4 and 3.5. Theorem 3.4 established the convergence of the averaged plans $\bar{\mathbf{p}}_{[T]} = \frac{1}{T} \sum_{t=1}^T \mathbf{p}_{(t)}$ and $\bar{\mathbf{a}}_{[T]} = \frac{1}{T} \sum_{t=1}^T \mathbf{a}_{(t)}$ to a set $K_P \times K_A$ of feasible saddle points of the overall social player cost function Ω - an insight that served as a linchpin for the ensuing theoretical argumentation regarding properties of the negotiation setup. Both theorems enabled us to establish the connection of the original planning problem to the planning outcome spawned by our mechanism and allowed us to infer social optimality results.

For the sake of a coherent exposition we decided to omit the respective proofs in Section 3.2.2. This gap shall be closed now.

C.1 Maxmin and minmax inequalities

Before commencing with the proofs we need to establish some preliminary results. As always, we assume that all agents employ no-regret learning and thus, the synthesized agents P and A have sublinear regret bounds Δ_P and Δ_A , respectively.

Lemma C.1. *Let $\bar{\mathbf{p}}_{[T]} := \frac{1}{T} \sum_{t=1}^T \mathbf{p}_{(t)}$, $\bar{\mathbf{a}}_{[T]} := \frac{1}{T} \sum_{t=1}^T \mathbf{a}_{(t)}$ be the respective average strategies until time step T .*

We have:

$$\min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) \leq \max_{\mathbf{a}} \Omega(\bar{\mathbf{p}}_{[T]}, \mathbf{a}) \leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T}.$$

and

$$\min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) \leq \min_{\mathbf{p}} \Omega(\mathbf{p}, \bar{\mathbf{a}}_{[T]}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T} \leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T}.$$

Proof. $\min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) \leq \max_{\mathbf{a}} \Omega(\bar{\mathbf{p}}_{[T]}, \mathbf{a})$

$$= \max_{\mathbf{a}} \Omega\left(\frac{1}{T} \sum_{t=1}^T \mathbf{p}_{(t)}, \mathbf{a}\right)$$

$$\leq^8 \max_{\mathbf{a}} \frac{1}{T} \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a})$$

⁸ $\Omega(\cdot, \mathbf{a})$ convex $\forall \mathbf{a}$.

$$\begin{aligned}
&\leq^9 \frac{1}{T} \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)}) + \frac{\Delta_A(T)}{T} \\
&\leq^{10} \min_{\mathbf{p}} \frac{1}{T} \sum_{t=1}^T \Omega(\mathbf{p}, \mathbf{a}_{(t)}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T} \\
&\leq^{11} \min_{\mathbf{p}} \Omega(\mathbf{p}, \bar{\mathbf{a}}_{[T]}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T} \\
&\leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T}
\end{aligned}$$

□

Remark C.2. The proof of the above Lemma was inspired by Freund and Shapire's alternative proof for *von Neumann's minimax theorem* which they presented in [FS96]. They considered a zero-sum matrix game between a no-regret learning player and a best-response environment (adversary). Their payoff-function model allowed them to leverage bi-linearity in their derivations. In contrast, we consider both player and adversary to be no-regret learners and assume player cost Ω to be merely continuous and convex-concave. (In addition, we tolerate the (synthesized) adversary A to have a slightly different revenue function $\Omega_A(\mathbf{p}, \mathbf{a}) = \Omega(\mathbf{p}, \mathbf{a}) - \langle \mathbf{c}, \mathbf{p} \rangle$).¹²

Remark C.3. It is standard knowledge that for any real-valued function κ on a nonempty product set $X \times Y$, we have $\sup_{\mathbf{x} \in X} \inf_{\mathbf{y} \in Y} \kappa(\mathbf{x}, \mathbf{y}) \leq \inf_{\mathbf{y} \in Y} \sup_{\mathbf{x} \in X} \kappa(\mathbf{x}, \mathbf{y})$ [Roc70]. For our case, this fact translates to $\max_{\mathbf{a} \in F_A} \min_{\mathbf{p} \in F_P} \Omega(\mathbf{p}, \mathbf{a}) \leq \min_{\mathbf{p} \in F_P} \max_{\mathbf{a} \in F_A} \Omega(\mathbf{p}, \mathbf{a})$, since we work with compact sets and continuous functions.

We now have the means to show that the minmax- and the maxmin value coincide. Such a result was first proven in von Neumann's well-known minimax theorem ([vN28]) for the case of two-player, zero-sum matrix games (Theorem C.4).

Theorem C.4. *Let $\Omega : F_P \times F_A \rightarrow \mathbb{R}$ be the payoff function of the two- player game between synthesized player P and adversary A (cf. Sec. 3.2).*

Then we have: $\max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) = \min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a})$.

Proof. By construction of Ω we have: $\forall \mathbf{a} \in F_A : \Omega(\cdot, \mathbf{a}) : F_P \rightarrow \mathbb{R}$ convex and $\forall \mathbf{p} \in F_P : \Omega(\mathbf{p}, \cdot) : F_A \rightarrow \mathbb{R}$ concave.

We can learn to play the game employing our no-regret learning mechanism (with one synthesized player agent P and one synthesized adversarial agent A). We can then apply Lemma C.1 to obtain the inequality $\min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) \leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T}$. Letting $T \rightarrow \infty$ we can conclude $\min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) \leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a})$. From Lemma C.3 we also know $\max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) \leq \min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a})$. □

C.2 Proof of Theorem 3.5

With the help of the inequalities found above, we are able to prove Theorem 3.5. Remember, it stated that

$$\frac{1}{T} \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)}) \rightarrow \min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) \text{ (as } T \rightarrow \infty \text{)}.$$

⁹Lemma 3.3.

¹⁰Lemma 3.2.

¹¹ $\Omega(\mathbf{p}, \cdot)$ concave $\forall \mathbf{p}$.

¹²Remember, the actual payoff of the entire collective of adversarial agents is $\Omega_A(\mathbf{p}, \mathbf{a}) - \sum_{j=1}^n \beta_j(\nu_j(\mathbf{p}))$. However, since the player penalty $\beta_j(\nu_j(\mathbf{p}))$ in nature for overconsumption of resource j is not controllable by adversarial agents we could construe the game to be played by an adversary whose payoff function was given by Ω_A .

if P and A incur no regret.

Proof. From the inequalities in the proof of Lemma C.1 we know:

$$\min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a}) \leq \frac{1}{T} \sum_{t=1}^T \Omega(\mathbf{p}_{(t)}, \mathbf{a}_{(t)}) + \frac{\Delta_A(T)}{T} \leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T}.$$

Considering $\Delta_P, \Delta_A \in o(T)$ and the result from Theorem C.4 completes the proof. $\square$

C.3 Proof of Theorem 3.4

Theorem 3.4 played a central role in our theoretical considerations. Remember, it stated the following:

If synthesized player P and synthesized adversary A suffer sublinear external regret, the average strategies $\bar{\mathbf{p}}_{[T]} = \frac{1}{T} \sum_{t=1}^T \mathbf{p}_{(t)}$ of P and $\bar{\mathbf{a}}_{[T]} = \frac{1}{T} \sum_{t=1}^T \mathbf{a}_{(t)}$ of A converge to a subset of saddle-points of objective function Ω (as $T \rightarrow \infty$).

Proof. Remember, $F_P \subset V, F_A \subset \mathbb{R}^n$ denote the compact, convex feasible sets of the synthesized player P and adversary A , respectively. Before proceeding, we establish some notation. Let $v := \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a})$. From Theorem C.4 we know that also $v = \min_{\mathbf{p}} \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a})$. Let the set of all saddle-points of Ω be denoted by E .

$E = \{(\mathbf{p}, \mathbf{a}) \in F_P \times F_A \mid \forall \tilde{\mathbf{p}} \in F_P \forall \tilde{\mathbf{a}} \in F_A : \Omega(\mathbf{p}, \tilde{\mathbf{a}}) \leq \Omega(\mathbf{p}, \mathbf{a}) \leq \Omega(\tilde{\mathbf{p}}, \mathbf{a})\}$ is the set of all feasible saddle points of Ω (cf. [Roc70]) which correspond to the minimax-equilibria of the game.

Let $K_P := \{\mathbf{p} \in V \mid \forall e > 0, t \in \mathbb{N} \exists T > t : \|\mathbf{p} - \bar{\mathbf{p}}_{[T]}\| < e\}$ be the set of all cluster points of sequence $(\bar{\mathbf{p}}_{[T]})_{T \in \mathbb{N}}$, $K_A := \{\mathbf{a} \in \mathbb{R}^n \mid \forall e > 0, t \in \mathbb{N} \exists T > t : \|\mathbf{a} - \bar{\mathbf{a}}_{[T]}\| < e\}$ the set of all cluster points of sequence $(\bar{\mathbf{a}}_{[T]})_{T \in \mathbb{N}}$. We could restate the definition and say K_P consists of all vectors being the limit of a subsequence $(\mathbf{s}_T)_{T \in \mathbb{N}}$ of $(\bar{\mathbf{p}}_{[T]})_{T \in \mathbb{N}}$. The same way, K_A consists of all vectors that are limit of a subsequence $(\mathbf{u}_T)_{T \in \mathbb{N}}$ of $(\bar{\mathbf{a}}_{[T]})_{T \in \mathbb{N}}$. Since we operate in compact subsets of Hilbert-spaces we know $K_P \subset F_P, K_A \subset F_A$ and $K_A, K_P \neq \emptyset$ if $F_P, F_A \neq \emptyset$.

The sets K_P, K_A are of great interest because the sequences $(\bar{\mathbf{p}}_{[T]})_{T \in \mathbb{N}}, (\bar{\mathbf{a}}_{[T]})_{T \in \mathbb{N}}$ converge towards them, respectively. With abuse of notation, we can write $\lim_{T \rightarrow \infty} \bar{\mathbf{p}}_{[T]} = K_P, \lim_{T \rightarrow \infty} \bar{\mathbf{a}}_{[T]} = K_A$.

In order to prove the claim of this theorem, we will show $K_P \times K_A \subset E$, i.e. $(\tilde{\mathbf{p}}, \tilde{\mathbf{a}}) \in E, \forall \tilde{\mathbf{p}} \in K_P, \tilde{\mathbf{a}} \in K_A$.

To do so, it suffices to show

$$\forall \tilde{\mathbf{p}} \in K_P, \tilde{\mathbf{a}} \in K_A : \max_{\mathbf{a}} \Omega(\tilde{\mathbf{p}}, \mathbf{a}) \leq \Omega(\tilde{\mathbf{p}}, \tilde{\mathbf{a}}) \leq \min_{\mathbf{p}} \Omega(\mathbf{p}, \tilde{\mathbf{a}}).$$

We introduce the well-defined functions: $\psi : F_A \rightarrow \mathbb{R}, \mathbf{a} \mapsto \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a})$ and $\phi : F_P \rightarrow \mathbb{R}, \mathbf{p} \mapsto \max_{\mathbf{a}} \Omega(\mathbf{p}, \mathbf{a})$ and note, $\forall \mathbf{a}, \mathbf{p} : \psi(\mathbf{a}) \leq \Omega(\mathbf{p}, \mathbf{a}) \leq \phi(\mathbf{p})$.

Combining the well-known maxmin inequality (cf. Remark C.3) and the results from Lemma C.1, we conclude for all $T \in \mathbb{N}$:

$$\begin{aligned}
& \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) \\
& \leq \max_{\mathbf{a}} \Omega(\bar{\mathbf{p}}_{[T]}, \mathbf{a}) \\
& \leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T} \\
& \text{and} \\
& \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) \\
& \leq \min_{\mathbf{p}} \Omega(\mathbf{p}, \bar{\mathbf{a}}_{[T]}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T} \\
& \leq \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) + \frac{\Delta_P(T)}{T} + \frac{\Delta_A(T)}{T}.
\end{aligned}$$

Hence, the following limits exist and are as follows:

$$\lim_{T \rightarrow \infty} \phi(\bar{\mathbf{p}}_{[T]}) = \lim_{T \rightarrow \infty} \max_{\mathbf{a}} \Omega(\bar{\mathbf{p}}_{[T]}, \mathbf{a}) = \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) = v \quad (10)$$

$$v = \max_{\mathbf{a}} \min_{\mathbf{p}} \Omega(\mathbf{p}, \mathbf{a}) = \lim_{T \rightarrow \infty} \min_{\mathbf{p}} \Omega(\mathbf{p}, \bar{\mathbf{a}}_{[T]}) = \lim_{T \rightarrow \infty} \psi(\bar{\mathbf{a}}_{[T]}). \quad (11)$$

Let $\tilde{\mathbf{p}} \in K_P, \tilde{\mathbf{a}} \in K_A$ be arbitrary choices. By definition of the cluster sets there exists a subsequence $(s_T)_T$ of $(\bar{\mathbf{p}}_{[T]})_T$ and a subsequence $(\mathbf{u}_T)_T$ of $(\bar{\mathbf{a}}_{[T]})_T$, such that : $s_T \rightarrow \tilde{\mathbf{p}} \wedge \mathbf{u}_T \rightarrow \tilde{\mathbf{a}}$ ($T \rightarrow \infty$). We know Ω, ϕ, ψ are continuous (cf. Sec. 3 and Appendix B, respectively). Hence, $v = \lim_{T \rightarrow \infty} \psi(\bar{\mathbf{a}}_{[T]}) = \lim_{T \rightarrow \infty} \psi(\mathbf{u}_{[T]}) = \psi(\tilde{\mathbf{a}})$. Analogously, we obtain: $v = \phi(\tilde{\mathbf{p}})$. Hence, $v = \phi(\tilde{\mathbf{p}}) \geq \Omega(\tilde{\mathbf{p}}, \mathbf{a}), \forall \mathbf{a}$ and thus, $v \geq \Omega(\tilde{\mathbf{p}}, \tilde{\mathbf{a}}) \geq \psi(\tilde{\mathbf{a}}) = v$. (Since $\tilde{\mathbf{p}}, \tilde{\mathbf{a}}$ were arbitrary choices we conclude $\Omega(K_P, K_A) = \{v\}$.)

With the help of Eq. 10 and Eq. 11, we can now conclude that $(\tilde{\mathbf{p}}, \tilde{\mathbf{a}})$ is a saddle point:

$$\max_{\mathbf{a}} \Omega(\tilde{\mathbf{p}}, \mathbf{a}) = \phi(\tilde{\mathbf{p}}) = v = \Omega(\tilde{\mathbf{p}}, \tilde{\mathbf{a}}) = v = \psi(\tilde{\mathbf{a}}) = \min_{\mathbf{p}} \Omega(\mathbf{p}, \tilde{\mathbf{a}}).$$

Since $\tilde{\mathbf{p}}, \tilde{\mathbf{a}}$ were arbitrary choices we conclude $K_P \times K_A \subset E$. $\square$

D Nonnegativity of remainder fee r_j

As always, let a^j be A_j 's plan and let $\mathbf{p}$ denote the overall player plan. In Sec. 3.1, we described how each player agent P_i transfers a payment $a^j \langle \mathbf{l}_{ji}, \mathbf{p}_i \rangle - \beta_{ji}(\mathbf{p}) - d_{ji} r_j(a^j)$ to each adversarial agent A_j . Furthermore, we mentioned that the magnitude of the remainder $r_j(a^j) = a^j y_j + \beta_j^*(a^j)$ was always nonnegative. We will proof this claim now. As a first step, we show that conjugation reverses inequalities:

Lemma D.1. *Let $d \in \mathbb{N}$, $\phi : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{-\infty, \infty\}$, $\gamma : \mathbb{R}^d \rightarrow \mathbb{R} \cup \{-\infty, \infty\}$ be convex functions. If $\forall \mathbf{x} \in \mathbb{R}^d : \phi(\mathbf{x}) \leq \gamma(\mathbf{x})$, then $\forall \mathbf{y} \in \mathbb{R}^d : \phi^*(\mathbf{y}) \geq \gamma^*(\mathbf{y})$.*

Proof. According to the premise we have $\forall \mathbf{x} : \phi(\mathbf{x}) \leq \gamma(\mathbf{x})$ and consequently, $\forall \mathbf{x} : -\phi(\mathbf{x}) \geq -\gamma(\mathbf{x})$. Hence, $\forall \mathbf{y} : \phi^*(\mathbf{y}) = \sup_{\mathbf{x}} \langle \mathbf{x}, \mathbf{y} \rangle - \phi(\mathbf{x}) \geq \sup_{\mathbf{x}} \langle \mathbf{x}, \mathbf{y} \rangle - \gamma(\mathbf{x}) = \gamma^*(\mathbf{y})$. $\square$

We can now show the desired result¹³:

Lemma D.2 (Nonnegativity of remainder r_j). *For all $a^j \geq 0$, we have $r_j(a^j) \geq 0$.*

¹³Remember, the feasible set was always contained in the positive half-space, i.e. $F_A \subset \mathbb{R}_+^n$. Therefore, we only need to consider nonnegative taxes a^j

Proof. Since $a^j, y_j \geq 0$ we only have to show that $\beta_j^*(a^j) \geq 0$. In Sec. 3.1, we required $\beta_j(x) \leq 0, \forall x \in \mathbb{R}$. Thus, $\beta_j(x) \leq \chi(x)$ where $\chi(x) = \begin{cases} 0, & x \leq 0 \\ \infty, & \text{otherwise} \end{cases}$ is the convex indicator function of the non-positive half-space.

According to Lemma D.1 we have $\forall y \in \mathbb{R} : \beta_j^*(y) \geq \chi^*(y)$. Referring to the definition of a convex function's conjugate it is easy to verify that $\chi^*(y) = \begin{cases} 0, & y \geq 0 \\ \infty, & \text{otherwise} \end{cases}$.

Hence, $\forall a^j \geq 0 : \beta_j^*(a^j) \geq 0$.

□

E Hard Inter-Agent Constraints

Throughout the paper we assumed that all inter-agent constraints (resource constraints) in nature were soft. That is, player agents can violate the constraints as much as they like - as long as they are willing to pay the increased price for such a behavior. In the light of our resource interpretation of these constraints this translates to the property of nature that resources are unlimited but tend to become arbitrarily expensive with increasing demand once a certain level of overall resource consumption is exceeded. As an example for a plausible domain where this model may be accurate, consider our network routing example where the resources are bandwidth-limitations on the communication links. In a conceivable scenario, we could imagine that once overall usage of a particular link threatens to exceed the maximal bandwidth, the network provider rents additional bandwidth from competitors who charge at expensive rates monotonically increasing with the rented bandwidth. Alternatively, the extra cost for resource constraint violation could encode the delay due to congestion.

However, one may wish to apply our method to planning in environments with hard constraints, i.e. in scenarios where a soft constraint model is not applicable but where the available resources are strictly limited.

In such settings the individual objective functions are linear and the descriptions of the feasible sets of each player agent contain the constraints $\{\nu_j(\mathbf{p}) \leq 0\}_{j=1\dots n}$ where $\nu_j(\mathbf{p}) = \langle \mathbf{l}_j, \mathbf{p} \rangle - y_j$ (cf. Eq. 1). In other words, instead of having $\min_{\mathbf{p}_i} \omega_{P_i}(\mathbf{p})$ s.t. : $\mathbf{p}_i \in F_{P_i}$ as its individual convex optimization problem, each player agent P_i should optimally solve:

$$\min_{\mathbf{p}_i} \langle \mathbf{c}_i, \mathbf{p}_i \rangle \text{ s.t. : } \mathbf{p}_i \in F_{P_i} \cap H_{\mathbf{p}_{\neg i}} \quad (12)$$

where $H_{\mathbf{p}_{\neg i}} := \{\mathbf{p}_i | \nu_1(\mathbf{p}) \leq 0, \dots, \nu_n(\mathbf{p}) \leq 0\}$. Note, subscript $\mathbf{p}_{\neg i}$ indicates $H_{\mathbf{p}_{\neg i}}$ is parameterized by $\mathbf{p}_{\neg i}$ being the overall plan of all players other than P_i . The socially optimal solution can be found as the solution to the optimization problem:

$$\min_{\mathbf{p}} \langle \mathbf{c}, \mathbf{p} \rangle \text{ s.t. : } \mathbf{p} \in F_p \cap H \quad (13)$$

where $H := \{\mathbf{p} | \nu_1(\mathbf{p}) \leq 0, \dots, \nu_n(\mathbf{p}) \leq 0\}$.

While our coordination mechanism needs to model the inter-agent constraints as being soft our mechanism is also suitable to achieve approximate coordination in the hard constraint scenario.

That is, it can be set up such that the overall player part $\tilde{\mathbf{p}}$ of a negotiation outcome is an approximate solution of optimization problem (13).

To achieve this, we simply need to define an artificial penalization function β_j for each hard inter-agent constraint $\nu_j(\mathbf{p}) \leq 0$ and assign an adversarial agent A_j to it as before ($j = 1, \dots, n$). While the true cost in nature is $\langle \mathbf{c}, \mathbf{p} \rangle$ we invoke the negotiation setup with the augmented social cost function $\omega(\mathbf{p}) = \langle \mathbf{c}, \mathbf{p} \rangle + \sum_j \beta_j(\nu_j(\mathbf{p}))$ and feasible set F_p . After learning in the repeated game is concluded, we know that the negotiation outcome $\tilde{\mathbf{p}}$ is optimal with respect to the regularized social cost function we defined, i.e. we have $\tilde{\mathbf{p}} = \operatorname{argmin}_{\mathbf{p} \in F_p} \omega(\mathbf{p}) = \operatorname{argmin}_{\mathbf{p} \in F_p} \langle \mathbf{c}, \mathbf{p} \rangle + \sum_j \beta_j(\nu_j(\mathbf{p}))$. In order to make sure that this solution is approximately optimal with respect to optimization problem (13) we need to assert that $\min_{\mathbf{p} \in F_p \cap H} \langle \mathbf{c}, \mathbf{p} \rangle \approx \langle \mathbf{c}, \tilde{\mathbf{p}} \rangle$ and that $\operatorname{dist}(\tilde{\mathbf{p}}, F_p \cap H) < \epsilon$ for some $\epsilon > 0$ which we predefine to encode the extend to which we consider constraint violations to be negligible.

Obviously, we can accomplish this by defining regularization functions β_j to penalize constraint violations accordingly. A simple yet perfectly suitable choice for β_j may be the prominent *hinge-loss* function which satisfies all the required model assumptions such as continuity, convexity and monotonicity. That is we define $\beta_j(\nu_j(\mathbf{p})) = \max(0, u_j \nu_j(\mathbf{p}))$ where u_j is a parameter defining the slope determining how much each unit of constraint violation is penalized.



MACHINE LEARNING
DEPARTMENT

Carnegie Mellon University
5000 Forbes Avenue
Pittsburgh, PA 15213

Carnegie Mellon.

Carnegie Mellon University does not discriminate and Carnegie Mellon University is required not to discriminate in admission, employment, or administration of its programs or activities on the basis of race, color, national origin, sex or handicap in violation of Title VI of the Civil Rights Act of 1964, Title IX of the Educational Amendments of 1972 and Section 504 of the Rehabilitation Act of 1973 or other federal, state, or local laws or executive orders.

In addition, Carnegie Mellon University does not discriminate in admission, employment or administration of its programs on the basis of religion, creed, ancestry, belief, age, veteran status, sexual orientation or in violation of federal, state, or local laws or executive orders. However, in the judgment of the Carnegie Mellon Human Relations Commission, the Department of Defense policy of, "Don't ask, don't tell, don't pursue," excludes openly gay, lesbian and bisexual students from receiving ROTC scholarships or serving in the military. Nevertheless, all ROTC classes at Carnegie Mellon University are available to all students.

Inquiries concerning application of these statements should be directed to the Provost, Carnegie Mellon University, 5000 Forbes Avenue, Pittsburgh PA 15213, telephone (412) 268-6684 or the Vice President for Enrollment, Carnegie Mellon University, 5000 Forbes Avenue, Pittsburgh PA 15213, telephone (412) 268-2056

Obtain general information about Carnegie Mellon University by calling (412) 268-2000