

Multiplayer Performative Prediction: Learning in Decision-Dependent Games

Adhyayan Narang

ADHYAN@UW.EDU

*Department of Electrical and Computer Engineering
University of Washington
Seattle, WA 98195-4322, USA*

Evan Faulkner

EVANJF5@UW.EDU

*Department of Electrical and Computer Engineering
University of Washington
Seattle, WA 98195-4322, USA*

Dmitriy Drusvyatskiy

DDRUSV@UW.EDU

*Department of Mathematics
University of Washington
Seattle, WA 98195-4322, USA*

Maryam Fazel

MFAZEL@UW.EDU

*Department of Electrical and Computer Engineering
University of Washington
Seattle, WA 98195-4322, USA*

Lillian J. Ratliff

RATLIFFL@UW.EDU

*Department of Electrical and Computer Engineering
University of Washington
Seattle, WA 98195-4322, USA*

Editor: Francesco Orabona

Abstract

Learning problems commonly exhibit an interesting feedback mechanism wherein the population data reacts to competing decision makers' actions. This paper formulates a new game theoretic framework for this phenomenon, called *multi-player performative prediction*. We focus on two distinct solution concepts, namely (i) performatively stable equilibria and (ii) Nash equilibria of the game. The latter equilibria are arguably more informative, but are generally computationally difficult to find since they are solutions of nonmonotone games. We show that under mild assumptions, the performatively stable equilibria can be found efficiently by a variety of algorithms, including repeated retraining and the repeated (stochastic) gradient method. We then establish transparent sufficient conditions for strong monotonicity of the game and use them to develop algorithms for finding Nash equilibria. We investigate derivative free methods and adaptive gradient algorithms wherein each player alternates between learning a parametric description of their distribution and gradient steps on the empirical risk. Synthetic and semi-synthetic numerical experiments illustrate the results.

Keywords: performative prediction, stochastic games, stochastic optimization, distributional shift, stochastic gradient method.

1. Introduction

Supervised learning theory and algorithms crucially rely on the training and testing data being generated from the same distribution. This assumption, however, is often violated in contemporary applications because data distributions may “shift” in reaction to the decision maker’s actions. Indeed, supervised learning algorithms are increasingly being trained on data that is generated by strategic or even adversarial agents, and deployed in environments that react to the decisions that the algorithm makes. In such settings, the model learned on the training data may be unsuitable for downstream inference and prediction tasks.

The method most commonly used in machine learning practice to address such distributional shifts is to periodically retrain the model to adapt to the changing distribution (Diethe et al., 2019; Wu et al., 2020). Consequently, it is important to understand when such retraining heuristics converge and what types of solutions they find. Despite the ubiquity of retraining heuristics in practice, one should be aware that training without consideration of strategic effects or decision-dependence can lead to unintended consequences, including reinforcing bias. This is a concern for applications with potentially significant social impact, such as predictive policing (Lum and Isaac, 2016), criminal sentencing (Angwin et al., 2016; Courtland, 2018), pricing equity in ride-share markets (Chen et al., 2015), and loan or job procurement (Bartlett et al., 2019).

Optimization over decision-dependent probabilities has classical roots in operations research; see for example the review article of (Hellemo et al., 2018) and references therein. The more recent work of (Perdomo et al., 2020), motivated by the strategic classification literature (Dong et al., 2018; Hardt et al., 2016; Miller et al., 2020), sets forth an elegant framework—aptly named *performative prediction*—for modeling decision-dependent data distributions in machine learning settings. There is now extensive research that develops algorithms for performative prediction by leveraging advances in convex optimization (Drusvyatskiy and Xiao, 2023; Miller et al., 2021; Mendler-Dünner et al., 2020; Perdomo et al., 2020; Brown et al., 2022).

The existing strategic classification and performative prediction literature focuses solely on the interplay between a single learner and the population that reacts to the learner’s actions. However, learning algorithms in practice are often deployed alongside other algorithms which may even be competing with one another. Concrete examples to keep in mind are those of college admissions and loan procurement, wherein the applicants may tailor their profile to make them more desirable for the college of their choice, or the loan with the terms (such as interest rate) that match the applicant’s current socio-economic and fiscal situation. In these cases, there are multiple competing learners (colleges, banks) and the population reacts based on the admissions policies of all the colleges (or banks) simultaneously. Examples of this type are widespread in applications; we provide further motivating vignettes in Section 3.

1.1 Contributions

We formulate the first game theoretic model for decision-dependent learning in the presence of competition, called *multi-player performative prediction*.¹ This is a new class of stochastic games that model a variety of machine learning problems arising in many practical appli-

1. A preliminary version of this paper appeared in the Proceedings of the 25th International Conference on Artificial Intelligence and Statistics, 2022.

cations. The model captures, as a special case, important problems including strategic classification in settings with multiple decision-making entities that model learning when each entity’s data distribution depends on the action taken. It defines an entire new class of problems that can be studied to determine consequences, unintended or otherwise, of using machine learning algorithms (including classifiers and predictors) in settings where the “data” generated for training is produced by strategic users that react based on their own internal preferences. Indeed, strategic behavior can lead to feedback loops in the data which if not monitored can, e.g., reinforce institutional biases as we have seen in predictive policing (Ensign et al., 2018; Lum and Isaac, 2016). Such consequences may be further exacerbated or alleviated in competitive settings where populations of users are choosing amongst multiple firms providing different service qualities (Dean et al., 2022).

We focus on two solution concepts for such games: (i) performatively stable equilibria and (ii) Nash equilibria. The former arises naturally when decision-makers employ naïve repeated retraining algorithms. This is very common practice, and hence it is important to understand the equilibrium to which such algorithms converge and precisely when they do so. We show that performatively stable equilibria are sure to exist and to be unique under reasonable smoothness, convexity, and Lipschitz assumptions (Section 4). Moreover, repeated training and (stochastic) gradient methods succeed at finding such equilibrium strategies. The finite time efficiency estimates (or iteration complexity) we obtain reduce to state-of-the art guarantees in the single player setting.

In some applications, a performatively stable equilibrium may be a poor solution concept, and instead a Nash equilibrium may be desirable; indeed, the latter is game theoretically meaningful in the sense that players have no incentive to change their action even when taking into consideration the decision-dependence in the data distribution and how this impacts the expected loss experienced by the decision maker. The concept of a performatively stable equilibrium does not have this feature, meaning that decision-makers may have a direction in which they can adjust their strategy and improve their loss. In particular, as machine learning algorithms become more sophisticated, in the sense that at the time of learning decision-dependence is taken into consideration, a more natural equilibrium concept is a Nash equilibrium. Aiming towards algorithms for finding Nash equilibria, we develop transparent conditions ensuring strong monotonicity of the game (Section 5). Assuming that the game is strongly monotone, we then discuss a number of algorithms that take into consideration different information structures available to the players for finding Nash equilibria (Section 6). In particular, derivative-free methods are immediately applicable but have a high sample complexity $\mathcal{O}(\frac{d^2}{\epsilon^2})$ (Bravo et al., 2018; Drusvyatskiy et al., 2022). Seeking faster algorithms, we introduce an additional assumption that the data distribution depends linearly on the performative effects of all the players. When the players know explicitly how the distribution depends on their own performative effects, but not those of their competitors, a simple stochastic gradient method is applicable and comes equipped with an efficiency guarantee of $\mathcal{O}(\frac{d}{\epsilon})$. Allowing players to know their own performative effects may be unrealistic in some settings. Consequently, we propose an *adaptive* algorithm in the setting when the data distribution has an amenable parametric description. In the algorithm, the players alternate between estimating the parameters of the distribution and optimizing their loss, again with only empirical samples of their individual gradients given

the estimated parameters. The sample complexity for this algorithm, up to variance terms, matches the rate $\mathcal{O}(\frac{d}{\epsilon})$ of the stochastic gradient method.

Finally, we present illustrative numerical experiments using both a synthetic example to validate the theoretical bounds, and a semi-synthetic example generated using data from multiple ride-share platforms (Section 7). Further experiments are contained in Appendix E.

1.2 Related Work

Performative Prediction. The multiplayer setting in the present paper is inspired by the single player performative prediction framework introduced by (Perdomo et al., 2020), and further refined by (Mendler-Dünner et al., 2020) and (Miller et al., 2021). These works introduce the two distinct concepts of (1) performative optimality and (2) performative stability. The former are optimal points of the optimization problem induced by decision-dependent data distributions, whereas the latter are fixed points of the repeated retraining procedure. Subsequently, (Drusvyatskiy and Xiao, 2023) showed that a variety of popular gradient-based algorithms—also seeking performatively stable points—in the decision-dependent setting can be understood as the analogous algorithms applied to a certain static problem corrupted by a vanishing bias. In general, however, performative stability does not imply performative optimality. Seeking to develop algorithms for finding performatively optimal points, (Miller et al., 2021) provide sufficient conditions for the prediction problem to be convex. Drawing connections with the present paper, observe that in the trivial “game” with a single decision-maker, the notion of a Nash equilibrium trivially reduces to a performatively optimal point.

For decision-dependent distributions described as location families, (Miller et al., 2021) additionally introduce a two-stage algorithm for finding performatively optimal points. The paper (Izzo et al., 2022) instead focuses on algorithms that estimate gradients with finite differences. The performative prediction framework is largely motivated by the problem of strategic classification (Hardt et al., 2016). This problem has been studied extensively from the perspective of causal inference (Bechavod et al., 2020; Miller et al., 2020) and convex optimization (Dong et al., 2018).

Another line of work in performative prediction has focused on the setting in which the environment evolves dynamically in time or experiences time drift. This line of work is more closely related to reinforcement learning wherein a decision maker attempts to maximize their reward over time given that the stochastic environment depends on their decision. In particular, (Brown et al., 2022) formulate a time-dependent performative prediction problem such that the decision-maker seeks to optimize the stationary reward—i.e., the reward under the fixed point distribution induced by the player’s decision. Repeated retraining algorithms seeking the performatively stable solution to this problem are studied. In contrast, (Ray et al., 2022) study performative prediction in geometrically decaying environments, and provide conditions and algorithms that lead to the performatively optimal solution. The papers (Cutler et al., 2021; Li and Wai, 2022; Wood et al., 2021) study performative prediction problems wherein the environment is drifting not only due to the action of the decision maker but also in time. These two papers analyze the tracking efficiency of the proximal stochastic gradient method and projected gradient descent under time drift.

Gradient-Based Learning in Continuous Games. There is a broad and growing literature on learning in games. We focus on the most relevant subset: gradient-based learning in continuous games. In his seminal work, (Rosen, 1965) showed that convex games which are *diagonal strictly convex* admit a unique Nash equilibrium and that gradient play converges to it. There is a large literature extending this work to more general games. For instance, (Ratliff et al., 2016) provide a characterization of Nash equilibria in non-convex continuous games, and show that continuous time gradient dynamics locally converge to Nash; building on this work, (Chasnov et al., 2020a) provide local convergence rates that extend to global rates when the game admits a potential function or is strongly monotone.

Under the assumption of strong monotonicity, the iteration complexity of stochastic and derivative-free gradient methods has also been obtained (Mertikopoulos and Zhou, 2019; Bravo et al., 2018; Drusvyatskiy et al., 2022). Relaxing strong monotonicity to monotonicity, Tatarenko and Kamgarpour (2019, 2020) show that the stochastic gradient and derivative free gradient methods—i.e., where players use a single-point query of the loss to construct an estimate of their individual gradient of a smoothed version of their loss function—converge asymptotically. The approach to deal with the lack of strong monotonicity is to add a regularization term that decays to zero asymptotically. The update players employ in this regularized game is then analyzed as a stochastic gradient method with an additional bias term. We take a similar perspective to Tatarenko and Kamgarpour (2019, 2020) and (Drusvyatskiy and Xiao, 2023) in the analysis of all the algorithms we study—namely, we view the updates as a stochastic gradient method with additional bias.

Stochastic programming. Stochastic optimization problems with decision-dependent uncertainties have appeared in the classical stochastic programming literature, such as (Ahmed, 2000; Dupacová, 2006; Jonsbråten et al., 1998; Rubinstein and Shapiro, 1993; Varaiya and Wets, 1988). We refer the reader to the recent paper (Hellemo et al., 2018), which discusses taxonomy and various models of decision dependent uncertainties. An important theme of these works is to utilize structural assumptions on how the decision variables impact the distributions. Consequently, these works sharply deviate from the framework explored in (Perdomo et al., 2020; Mendler-Dünner et al., 2020; Miller et al., 2021) and from our paper. Time-varying problems have also been studied under the title “nonstationary optimization problems” in, e.g., (Gaivoronskii, 1978; Ermoliev, 1988), where it is assumed that the time varying functions converges to a limit but there is no explicit assumption on decision or state-feedback.

2. Notation and Preliminaries

This section records basic notation that we will use. A reader that is familiar with convex games and the Wasserstein-1 distance between probability measures may safely skip this section. Throughout, we let $\mathbb{R}^d$ denote a d -dimensional Euclidean space, with inner product $\langle \cdot, \cdot \rangle$ and the induced norm $\|x\| = \sqrt{\langle x, x \rangle}$. The projection of a point $y \in \mathbb{R}^d$ onto a set $\mathcal{X} \subset \mathbb{R}^d$ is denoted

$$\text{proj}_{\mathcal{X}}(y) = \underset{x \in \mathcal{X}}{\operatorname{argmin}} \|x - y\|.$$

The normal cone to a convex set $\mathcal{X}$ at $x \in \mathcal{X}$, denoted by $N_{\mathcal{X}}(x)$, is the set

$$N_{\mathcal{X}}(x) = \{v \in \mathbb{R}^d : \langle v, y - x \rangle \leq 0 \quad \forall y \in \mathcal{X}\}.$$

We say that f is C^1 -smooth if the mapping $x \mapsto \nabla_x f(x)$ is well-defined and continuous.

2.1 Convex Games and Monotonicity

Fix an index set $[n] = \{1, \dots, n\}$, integers d_i for $i \in [n]$, and set $d = \sum_{i=1}^n d_i$. Throughout, we decompose vectors $x \in \mathbb{R}^d$ as $x = (x_1, \dots, x_n)$ with $x_i \in \mathbb{R}^{d_i}$. Given an index i , we abuse notation and write $x = (x_i, x_{-i})$, where x_{-i} denotes the vector of all coordinates except x_i . A collection of functions $\mathcal{L}_i : \mathbb{R}^d \rightarrow \mathbb{R}$ and sets $\mathcal{X}_i \subset \mathbb{R}^{d_i}$, for $i \in [n]$, induces a game between n players, wherein each player i seeks to solve the problem

$$\min_{x_i \in \mathcal{X}_i} \mathcal{L}_i(x_i, x_{-i}). \quad (1)$$

Define the joint action space $\mathcal{X} = \mathcal{X}_1 \times \dots \times \mathcal{X}_n$. A vector $x^* \in \mathbb{R}^d$ is called a *Nash equilibrium* of the game (1) if the condition

$$x_i^* \in \operatorname{argmin}_{x_i \in \mathcal{X}_i} \mathcal{L}_i(x_i, x_{-i}^*) \quad \text{holds for each } i \in [n]. \quad (2)$$

Thus x^* is a Nash equilibrium if each player i has no incentive to deviate from x_i^* when the strategies of all other players remain fixed at x_{-i}^* .

The gradient of $\mathcal{L}_i(\cdot, \cdot)$ with respect to the argument x_i is denoted $\nabla_i \mathcal{L}_i(x_i, x_{-i}) = (\nabla_{x_{ij}} \mathcal{L}_i(x_i, x_{-i}))_{j=1}^{d_i} \in \mathbb{R}^{d_i}$. With this notation, we define the vector of individual gradients

$$H(x) := (\nabla_1 \mathcal{L}_1(x), \dots, \nabla_n \mathcal{L}_n(x)),$$

such that $H(x) \in \mathbb{R}^d$. This map arises naturally from writing down first-order optimality conditions corresponding to (1) for each player. Namely, we say that (1) is a C^1 -smooth convex game if the sets $\mathcal{X}_i$ are closed and convex, the functions $\mathcal{L}_i(\cdot, x_{-i})$ are convex (i.e., $\mathcal{L}_i$ is convex in x_i when x_{-i} are fixed), and the partial gradients $\nabla_i \mathcal{L}_i(x)$ exist and are continuous. Thus, the Nash equilibria x^* are characterized by the inclusion

$$0 \in H(x^*) + N_{\mathcal{X}}(x^*).$$

A C^1 -smooth convex game is called α -strongly monotone (for $\alpha \geq 0$) if it satisfies

$$\langle H(x) - H(x'), x - x' \rangle \geq \alpha \|x - x'\|^2 \quad \text{for all } x, x' \in \mathbb{R}^d.$$

If this condition holds with $\alpha = 0$, the game is simply called *monotone*. It is well-known from (Rosen, 1965) that α -strongly monotone games (with $\alpha > 0$) over convex, closed and bounded strategy sets $\mathcal{X}$ admit a *unique* Nash equilibrium x^* , which satisfies

$$\langle H(x), x - x^* \rangle \geq \alpha \|x - x^*\|^2 \quad \text{for all } x \in \mathcal{X}.$$

2.2 Probability Measures and Gradient Deviation

To simplify notation, we will always assume that when taking expectations with respect to a measure that the expectation exists and that integration and differentiation operations may be swapped whenever we encounter them. These assumptions are completely standard to justify under uniform integrability conditions.

We are interested in random variables taking values in a metric space. Given a metric space $\mathcal{Z}$ with metric $d(\cdot, \cdot)$, the symbol $\mathbb{P}(\mathcal{Z})$ will denote the set of Radon probability measures μ on $\mathcal{Z}$ with a finite first moment $\mathbb{E}_{z \sim \mu}[d(z, z_0)] < \infty$ for some $z_0 \in \mathcal{Z}$. We measure the deviation between two measures $\mu, \nu \in \mathbb{P}(\mathcal{Z})$ using the Wasserstein-1 distance:

$$W_1(\mu, \nu) = \sup_{h \in \text{Lip}_1} \left\{ \mathbb{E}_{X \sim \mu} [h(X)] - \mathbb{E}_{Y \sim \nu} [h(Y)] \right\}, \quad (3)$$

where Lip_1 denotes the set of 1-Lipschitz continuous functions $h: \mathcal{Z} \rightarrow \mathbb{R}$. Fix a function $g: \mathbb{R}^d \times \mathcal{Z} \rightarrow \mathbb{R}$ and a measure $\mu \in \mathbb{P}(\mathcal{Z})$, and define the expected loss

$$f_\mu(x) = \mathbb{E}_{z \sim \mu} g(x, z).$$

Throughout, the symbol $\nabla f_\mu(x)$ denotes the gradient of $f_\mu(\cdot)$ with respect to its argument x . The following standard result shows that the Wasserstein-1 distance controls how the gradient $\nabla f_\mu(x)$ varies with respect to μ ; see, for example, Drusvyatskiy and Xiao (2023, Lemmas 1.1, 2.1) for a short proof.

Lemma 1 (Gradient deviation) *Fix a function $g: \mathbb{R}^d \times \mathcal{Z} \rightarrow \mathbb{R}$ such that $g(\cdot, z)$ is C^1 for all $z \in \mathcal{Z}$ and the map $z \mapsto \nabla_x g(x, z)$ is β -Lipschitz continuous for any $x \in \mathbb{R}^d$. Fix now any measures $\mu, \nu \in \mathbb{P}(\mathcal{Z})$ such that $g(x, \cdot)$ is both μ and ν integrable for all x . Then we may exchange differentiation and integration $\nabla f_\mu(x) = \mathbb{E}_\mu \nabla g(x, z)$ and the estimate holds:*

$$\sup_x \|\nabla f_\mu(x) - \nabla f_\nu(x)\| \leq \beta \cdot W_1(\mu, \nu).$$

3. Decision-Dependent Games

We model the problem of n decision-makers, each facing a decision-dependent learning problem, as an n -player game. Each player $i \in [n]$ seeks to solve the decision-dependent optimization problem

$$\min_{x_i \in \mathcal{X}_i} \mathcal{L}_i(x_i, x_{-i}) \quad \text{where} \quad \mathcal{L}_i(x) := \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} \ell_i(x, z_i). \quad (4)$$

Throughout, we suppose that each set $\mathcal{X}_i$ lies in the Euclidean space $\mathbb{R}^{d_i}$ and we set $d = \sum_{i=1}^n d_i$. The loss function for the i 'th player is denoted as $\ell_i: \mathbb{R}^d \times \mathcal{Z}_i \rightarrow \mathbb{R}$, where $\mathcal{Z}_i$ is some metric space and $\mathcal{D}_i(x) \in \mathcal{P}(\mathcal{Z}_i)$ is a probability measure that depends on the joint decision $x \in \mathcal{X}$ and the player $i \in [n]$. Observe that the random variable z_i in the objective function of player i is governed by the distribution $\mathcal{D}_i(x)$, which crucially depends on the strategies $x = (x_1, \dots, x_n)$ chosen by *all* players. This is worth emphasizing: the parameters chosen by one player have an influence on the data seen by all other players. This

is one of the critical ways in which the problems for the different players are strategically coupled. The other is directly through the loss function ℓ_i which also depends on the joint decision x . These two sources of strategic coupling are why the game theoretic abstraction naturally arises. It is worth keeping in mind that in most practical settings (see the upcoming Vignettes 1 and 2), the loss functions $\ell_i(x, z_i)$ depend only on x_i , that is $\ell_i(x, z_i) \equiv \ell_i(x_i, z_i)$. If this is the case, we will call the game *separable* (which refers to separable losses, not distributions). Thus, for separable games, the coupling among the players is due entirely to the distribution $\mathcal{D}_i(x)$ that depends on the actions of all the players.

Remark 2 The decision-dependence in the distribution map may encode the reaction of strategic users in a population to the announced joint decision x ; hence, in these cases there is also a “game” between the decision-makers and the strategic users in the environment—a game with a different interaction structure known as a Stackelberg game (Von Stackelberg, 2010). This level of strategic interaction between decision-maker and strategic user is abstracted away to an aggregate level in $\mathcal{D}_i(x_i, x_{-i})$. The game between a single decision maker and the strategic user population has been studied widely (cf. Section 1.2). We leave it to future work to investigate both layers of strategic interaction simultaneously.

It is assumed that each player observes the other players’ actions. This is a reasonable assumption in our setup: if the data population (e.g., strategic users) are able to respond to the players’ deployed decisions x_i , the other players must be able to respond to these decisions as well. In essence, these decisions are publicly announced. We note that in practice the players may have to learn competitors’ decisions and this might result in asynchronous information. We comment in Sections 4 and 6 where we analyze algorithms in different information settings on how asynchronous feedback can be captured in the algorithms without fundamentally changing the convergence rate. We leave the question of players estimating the actions of competitors, and the question of users estimating players decision rules to future work.

The following vignettes based on practical applications motivate different types of strategic coupling.

Vignette 1 (Multiplayer forecasting) Players have the same decision-dependent data distribution—namely, $\mathcal{D} \equiv \mathcal{D}_i \equiv \mathcal{D}_j$ for all $i, j \in [n]$. Multiple mapping applications forecast the travel time between different locations, yet the realized travel time is collectively influenced by all their forecasts. The decision-dependent players are the mapping applications (firms). The decision x_i each player makes is the rule for recommending routes. Users choose routes, which then impact the realized travel time $z_i \equiv z \sim \mathcal{D}$ on the m different road segments in the network observed by all firms.

Vignette 2 Players have different distributions—i.e., $\mathcal{D}_i \neq \mathcal{D}_j$ for all $i, j \in [n]$.

- (a) **Multiplayer Strategic Classification.** Multiple universities classify students as accepted or rejected using applicant data, where each applicant designs their application to fit desiderata of multiple universities. The data $z_i \sim \mathcal{D}_i(x)$ is an application that university i receives, and as a decision-dependent player, each university i designs a classification rule x_i to determine which applicants are accepted. The goal of a university is to accept qualified candidates, and different types of universities predominantly cater

to different populations (e.g., liberal arts versus science and engineering), yet students may apply to multiple programs across many universities thereby resulting in distinct distributions $\mathcal{D}_i$ that depend on the joint decision rule x .

- (b) **Revenue Maximization via Demand Forecasting.** In a ride-share market, multiple platforms forecast demand for rides (respectively, supply of drivers) at different locations in order to optimize their revenue by using the forecast to set prices. In most ride-share markets, drivers and passengers participate in multiple platforms by, e.g., “price shopping”. Hence, the forecasted demand $z_i \sim \mathcal{D}_i(x)$ for platform i depends on their own decision x_i as well as their competitors’ decisions x_{-i} .

Prior formulations of decision dependent learning do not readily extend to the settings described in the vignettes without a game theoretic model. There are two natural solution concepts for the game (4). The first is the classical notion of Nash equilibrium; we repeat the definition here for ease of reference.

Definition 3 (Nash equilibrium) A vector $x^* \in \mathcal{X}$ is a *Nash equilibrium* of the game (4) if

$$x_i^* \in \operatorname{argmin}_{x_i \in \mathcal{X}_i} \mathcal{L}_i(x_i, x_{-i}^*) \quad \text{holds for each } i \in [n].$$

The game (4) can easily fail to be monotone even if it is separable and the loss functions $\ell_i(\cdot, z)$ are strongly convex. In Section 5, we develop sufficient conditions for (strong) monotonicity and use them to analyze algorithms for finding Nash equilibria. Further we provide examples that illustrate the sufficient conditions. We note however that the sufficient conditions we develop, which are in line with those in the single player setting (Miller et al., 2021), are quite restrictive.

On the other hand, there is an alternative solution concept, which is more amenable to numerical methods, and reduces to performatively stable points of (Perdomo et al., 2020) in the single player setting. The idea is to decouple the effects of a decision x on the integrand $\ell(x, z)$ and on the distribution $\mathcal{D}(x)$. Setting notation, any vector $y \in \mathcal{X}$ induces a static game, $\mathcal{G}(y)$ (without performative effects), wherein the distribution for player i is fixed at $\mathcal{D}_i(y)$:

$$\mathcal{G}(y) := \left(\mathbb{E}_{z_1 \sim \mathcal{D}_1(y)} \ell_1(x_1, x_{-1}, z_1), \dots, \mathbb{E}_{z_n \sim \mathcal{D}_n(y)} \ell_n(x_n, x_{-n}, z_n) \right), \quad (5)$$

meaning each player i seeks to solve

$$\min_{x_i \in \mathcal{X}_i} \mathcal{L}_i^y(x_i, x_{-i}) \quad \text{where} \quad \mathcal{L}_i^y(x_i, x_{-i}) := \mathbb{E}_{z_i \sim \mathcal{D}_i(y)} \ell_i(x_i, x_{-i}, z_i).$$

Notice that this is a parametric family of games, indexed by $y \in \mathcal{X}$, in which players do not take into consideration the dependence of $\mathcal{D}_i$ on the their actions; we refer to this class of games as “static games” since the distribution is fixed at y .

Definition 4 (Performatively stable equilibria) A strategy vector $x^* \in \mathcal{X}$ is a *performatively stable equilibrium* of the game (4) if it is a Nash equilibrium of the static game $\mathcal{G}(x^*)$ (with game $\mathcal{G}(\cdot)$ as defined above).

The difference between performatively stable equilibria and Nash equilibria is that the governing distribution for player i is fixed at $\mathcal{D}_i(x^*)$. Performatively stable equilibria have a clear intuitive meaning: each player i has no incentive to deviate from x^* having access only to data drawn from $\mathcal{D}(x^*)$. Notice that if the game (4) is separable—the typical setting—the static game $\mathcal{G}(y)$ is entirely decoupled for any y in the sense that the problem that player i aims to solve depends only on x_i and not on x_{-i} . This enables a variety of single player optimization techniques to extend to the computation of performatively stable equilibria.

Nash and performatively stable equilibria are typically distinct, even in the single player setting as explained in (Perdomo et al., 2020). This distinction is worth highlighting. Under mild smoothness assumptions,² taking the derivative directly implies the following expression for the gradient of the objective of player i :

$$\nabla_i \mathcal{L}_i(x_i, x_{-i}) = \underbrace{\mathbb{E}_{z_i \sim \mathcal{D}_i(x)} [\nabla_i \ell_i(x_i, x_{-i}, z_i)]}_{P_i} + \underbrace{\nabla_{u_i} \left(\mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x_{-i})} [\ell_i(x_i, x_{-i}, z_i)] \right)}_{Q_i} \Big|_{u_i=x_i}, \quad (6)$$

where $\nabla_i \ell_i(x_i, x_{-i}, z_i) = (\frac{\partial}{\partial x_{i,1}} \ell_i(x_i, x_{-i}, z_i), \dots, \frac{\partial}{\partial x_{i,d_i}} \ell_i(x_i, x_{-i}, z_i))$ is the gradient of ℓ_i with respect to the x_i argument, and

$$\nabla_{u_i} \left(\mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x_{-i})} [\ell_i(x_i, x_{-i}, z_i)] \right) \Big|_{u_i=x_i} = \left(\frac{\partial}{\partial u_{i,j}} \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x_{-i})} [\ell_i(x_i, x_{-i}, z_i)] \Big|_{u_{i,j}=x_{i,j}} \right)_{j=1}^{d_i}.$$

If x is a Nash equilibrium of the game (4) and $\mathcal{X} = \mathbb{R}^d$, then equality $0 = \nabla_i \mathcal{L}_i(x_i, x_{-i}) = P_i + Q_i$ holds for all $i \in [n]$. On the other hand, provided the loss functions ℓ_i are convex, the joint action x is a performatively stable equilibrium precisely when $P_i \equiv 0$ for all $i \in [n]$. This clearly shows that the two solution concepts are typically distinct, since performative stability in essence ignores the term Q_i on the right side of (6). It is an open question as to how these equilibrium concepts compare in terms of their efficiency relative to the social optimum. We explore this empirically in Section 7.

Before proceeding, we introduce some convenient notation that we use throughout. Fix two vectors $x = (x_1, \dots, x_n) \in \mathcal{X}$ and $z = (z_1, \dots, z_n) \in \mathcal{Z}_1 \times \dots \times \mathcal{Z}_n$. We then set

$$g_i(x, z_i) := \nabla_i \ell_i(x, z_i) \quad \text{and} \quad g(x, z) := (g_1(x, z_1), \dots, g_n(x, z_n)).$$

Taking expectations define

$$G_{i,y}(x) := \mathbb{E}_{z_i \sim \mathcal{D}_i(y)} g_i(x, z_i) \quad \text{and} \quad G_y(x) := (G_{1,y}(x), \dots, G_{n,y}(x)), \quad (7)$$

so that $G_y(\cdot)$ is the vector of individual gradients corresponding to the game 5. Notice that we may write

$$G_y(x) := \mathbb{E}_{z \sim \mathcal{D}_\pi(y)} g(x, z)$$

where $\mathcal{D}_\pi(y) := \mathcal{D}_1(y) \times \dots \times \mathcal{D}_n(y)$ is the product measure.

In the rest of the paper we impose the following assumption that is in line with the performative prediction literature.

2. cf. Assumption 5 in Section 5 for the precise assumptions needed for the product rule to apply.

Assumption 1 (Convexity and smoothness) *There exist constants $\alpha > 0$ and $\beta_i > 0$ such that for each $i \in [n]$, the following hold:*

1. *Each loss $\ell_i(x_i, x_{-i}, z_i)$ is C^1 -smooth in x_i and the map $z_i \mapsto \nabla_{z_i} \ell_i(x, z_i)$ is β_i -Lipschitz continuous for any $x \in \mathcal{X}$.*
2. *For any $y \in \mathcal{X}$, the game $\mathcal{G}(y)$ is α -strongly monotone—i.e., the game $\mathcal{G}(y)$ induced by the fixed joint action y satisfies*

$$\alpha \cdot \|x - x'\|^2 \leq \langle G_y(x) - G_y(x'), x - x' \rangle \quad \forall x, x' \in \mathcal{X}.$$

It is worth noting that in the setting where the losses are separable, the game $\mathcal{G}(y)$ is α -strongly monotone as long as each expected loss $\mathbb{E}_{z \sim \mathcal{D}_i(y)} \ell_i(x_i, z_i)$ is α -strongly convex in x_i . Assumption 1 alone *does not* imply convexity of the objective functions $\mathcal{L}_i(x_i, x_{-i})$ in x_i nor monotonicity of the game (4) itself. Sufficient conditions for convexity and strong monotonicity of the game are given in Section 5.

Next, we require the distributions $\mathcal{D}_i(x)$ to vary in a Lipschitz way with respect to x .

Assumption 2 (Lipschitz distributions) *For each $i \in [n]$, there exists $\gamma_i > 0$ satisfying*

$$W_1(\mathcal{D}_i(x), \mathcal{D}_i(y)) \leq \gamma_i \cdot \|x - y\| \quad \text{for all } x, y \in \mathcal{X}.$$

In this case, we define the constant $\rho := \sqrt{\sum_{i=1}^n (\frac{\beta_i \gamma_i}{\alpha})^2}$.

Theorem 7 implies that the constant ρ is fundamentally important for algorithms, since it characterizes settings in which algorithms can be expected to work.

The following is a direct consequence of Lemma 1.

Lemma 5 (Deviation in the vector of individual gradients) *Suppose Assumptions 1 and 2 hold. Then for every $x, y, y' \in \mathcal{X}$ and index $i \in [n]$, the estimates hold:*

$$\begin{aligned} \|G_{i,y}(x) - G_{i,y'}(x)\| &\leq \beta_i \gamma_i \cdot \|y - y'\|, \\ \|G_y(x) - G_{y'}(x)\| &\leq \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2} \cdot \|y - y'\|. \end{aligned}$$

Proof Using Lemma 1 and the standing Assumptions 1 and 2 we compute

$$\begin{aligned} \|G_{i,y}(x) - G_{i,y'}(x)\| &= \left\| \mathbb{E}_{z_i \sim \mathcal{D}_i(y)} \nabla_{z_i} \ell_i(x, z_i) - \mathbb{E}_{z_i \sim \mathcal{D}_i(y')} \nabla_{z_i} \ell_i(x, z_i) \right\| \\ &\leq \beta_i \cdot W_1(\mathcal{D}_i(y), \mathcal{D}_i(y')) \\ &\leq \beta_i \gamma_i \cdot \|y - y'\|. \end{aligned}$$

Therefore, we deduce that

$$\|G_y(x) - G_{y'}(x)\|^2 = \sum_{i=1}^n \|G_{i,y}(x) - G_{i,y'}(x)\|^2 \leq \sum_{i=1}^n \beta_i^2 \gamma_i^2 \cdot \|y - y'\|^2.$$

The proof is complete. ■

We end the section by noting that in principle one could bound the distance between performatively stable equilibria and Nash equilibria as follows.

Lemma 6 (Performatively stable & Nash equilibrium deviation) *Suppose that Assumptions 1 and 2 hold and that we are in the regime $\rho < 1$. Moreover, suppose that the expression (6) is valid and the loss functions $\ell_i(\cdot, x_{-i}, z_i)$ are L_i -Lipschitz on $\mathcal{X}_i$. Let $\bar{x}$ and x^* be, respectively, a Nash equilibrium and performatively stable equilibrium. Then the estimate holds:*

$$\|\bar{x} - x^*\| \leq \frac{\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2}}{\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}.$$

Proof Using strong monotonicity, we compute

$$\begin{aligned} \alpha \|\bar{x} - x^*\|^2 &\leq \langle G_{x^*}(\bar{x}) - G_{x^*}(x^*), \bar{x} - x^* \rangle \\ &\leq \langle G_{x^*}(\bar{x}), \bar{x} - x^* \rangle \\ &\leq \langle G_{\bar{x}}(\bar{x}), \bar{x} - x^* \rangle + \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2} \cdot \|\bar{x} - x^*\|^2, \end{aligned} \tag{8}$$

where the last inequality follows from Lemma 5. Next recall that by definition of Nash equilibrium and the expression (6) we have

$$0 \in G_{\bar{x}}(\bar{x}) + [Q_i]_{i=1}^n + N_{\mathcal{X}}(\bar{x}), \tag{9}$$

where Q_i is evaluated at $\bar{x}$. Note that the Kantorovich-Rubenstein dual representation for W_1 distance directly implies $\|[Q_i]_{i=1}^n\| \leq \sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2}$. Therefore we deduce

$$\langle G_{\bar{x}}(\bar{x}), \bar{x} - x^* \rangle \leq \sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} \cdot \|\bar{x} - x^*\|.$$

Combining (8)-(9) and rearranging and dividing by $\|\bar{x} - x^*\|$ completes the proof. $\blacksquare$

4. Algorithms for Finding Performatively Stable Equilibria

The previous section isolated two solution concepts for decision dependent games: Nash equilibria and performatively stable equilibria. In this section, we discuss existence of the latter and algorithms for finding these. We discuss three algorithms: repeated retraining, the repeated gradient method, and the repeated stochastic gradient method. While the first two are largely conceptual, the repeated stochastic gradient method is entirely implementable.

4.1 Repeated Retraining

Observe that performatively stable equilibria of (4) are precisely the fixed points of the map

$$\text{Nash}(x) := \{x' \in \mathcal{X} : x' \text{ is a Nash equilibrium of the game } \mathcal{G}(x)\}.$$

A natural algorithm is therefore *repeated retraining*, which is simply the fixed point iteration

$$x^{t+1} = \text{Nash}(x^t). \tag{10}$$

In the single player settings, this algorithm was studied in Perdomo et al. (2020). Unrolling notation, given a current decision vector x^t , the updated decision vector x^{t+1} is the Nash equilibrium of the game wherein each player $i \in [n]$ seeks to solve

$$\min_{y_i \in \mathcal{X}_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(x^t)} \ell_i(y_i, y_{-i}, z_i). \quad (11)$$

It is important to notice that since x^t is fixed, the game (11) is strongly monotone in light of Assumption 1. Thus, in iteration t , the players play a Nash equilibrium (i.e., a best response) in this game induced by the prior joint decision x^t . Importantly, despite the fact that x^{t+1} is a Nash equilibrium of a game in iteration t , the collective decision x^{t+1} is **not** a Nash equilibrium for the multiplayer performative prediction problem (4). In fact, players have an incentive to change their action since it is possible that by changing x_i^t , the change it induces in the distribution $\mathcal{D}_i(\cdot)$ reduces their expected loss.

The following theorem shows that in the regime $\rho < 1$, the game (4) admits a unique performatively stable equilibrium and repeated retraining converges linearly.

Theorem 7 (Repeated retraining) *Suppose Assumptions 1-2 hold and that we are in the regime $\rho < 1$. Then the game (4) admits a unique performatively stable equilibrium x^* and the repeated retraining process (10) converges linearly:*

$$\|x^{t+1} - x^*\| \leq \rho \cdot \|x^t - x^*\| \quad \text{for all } t \geq 0.$$

Proof We will show that the map $\text{Nash}(\cdot)$ is Lipschitz continuous with parameter ρ . To this end, consider two points x and x' and set $y := \text{Nash}(x)$ and $y' := \text{Nash}(x')$. Note that first order optimality conditions for y and y' guarantee

$$\langle G_x(y), y - y' \rangle \leq 0 \quad \text{and} \quad \langle G_{x'}(y'), y' - y \rangle \leq 0.$$

Strong monotonicity of the game $\mathcal{G}(x)$ therefore ensures

$$\begin{aligned} \alpha \cdot \|y - y'\|^2 &\leq \langle G_x(y) - G_x(y'), y - y' \rangle \\ &\leq \langle G_{x'}(y') - G_x(y'), y - y' \rangle \\ &\leq \|G_{x'}(y') - G_x(y')\| \cdot \|y - y'\| \\ &\leq \sqrt{\sum_{i=1}^n \gamma_i^2 \beta_i^2} \cdot \|x - x'\| \cdot \|y - y'\|, \end{aligned}$$

where the last inequality follows from Lemma 5. Dividing through by $\|y - y'\|$ guarantees that $\text{Nash}(\cdot)$ is indeed a contraction with parameter ρ . The result follows immediately from the Banach fixed point theorem. $\blacksquare$

Theorem 7 shows that the interesting parameter regime is $\rho < 1$, since outside of this setting performative equilibria may even fail to exist. It is worth noting that when the game (4) is separable, each iteration of repeated retraining (10) becomes

$$\min_{y_i \in \mathcal{X}_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(x^t)} \ell_i(y_i, z_i). \quad (12)$$

That is, the optimization problems faced by the players are entirely independent of each other. In the separable case, the regime when repeated retraining succeeds can be slightly enlarged from $\rho < 1$ to $\sum_{i=1}^n \left(\frac{\beta_i \gamma_i}{\alpha_i} \right)^2 < 1$, where α_i is the strong convexity constant of the loss for player i . This is the content of the following theorem, whose proof is a slight modification of the proof of Theorem 7.

Theorem 8 (Improved contraction for separable games) *Suppose that the game (4) is separable and that each loss function $\mathcal{L}_i^y(x_i) = \mathbb{E}_{z_i \sim \mathcal{D}_i(y)} \ell_i(x_i, z_i)$ is α_i -strongly convex in x_i for every $y \in \mathcal{X}$. Suppose moreover that Assumptions 1 and 2 hold, and that we are in the regime $\sum_{i=1}^n \left(\frac{\beta_i \gamma_i}{\alpha_i} \right)^2 < 1$. The game (4) admits a unique performatively stable equilibrium x^* and the repeated retraining process converges linearly:*

$$\|x^{t+1} - x^*\| \leq \sqrt{\sum_{i=1}^n \left(\frac{\beta_i \gamma_i}{\alpha_i} \right)^2} \cdot \|x^t - x^*\| \quad \text{for all } t \geq 0.$$

Proof We show that the map $\text{Nash}(\cdot)$ is Lipschitz continuous with parameter $\sqrt{\sum_{i=1}^n \left(\frac{\beta_i \gamma_i}{\alpha_i} \right)^2}$. To this end, consider two points x and x' and set $y := \text{Nash}(x)$ and $y' := \text{Nash}(x')$. Note that first order optimality conditions for y and y' guarantee

$$\langle G_{i,x}(y), y_i - y'_i \rangle \leq 0 \quad \text{and} \quad \langle G_{i,x'}(y'), y'_i - y_i \rangle \leq 0 \quad \text{for all } i \in [n].$$

Set $v = (\alpha_1^{-1}, \dots, \alpha_n^{-1})$ and let $\odot$ denote the Hadamard product between two vectors. Strong convexity of the loss functions therefore ensures

$$\begin{aligned} \|y - y'\|^2 &\leq \sum_i \alpha_i^{-1} \langle G_{i,x}(y) - G_{i,x}(y'), y_i - y'_i \rangle \\ &\leq \sum_i \alpha_i^{-1} \langle G_{i,x'}(y') - G_{i,x}(y'), y_i - y'_i \rangle \\ &= \langle v \odot (G_{x'}(y') - G_x(y')), y - y' \rangle \\ &\leq \|v \odot (G_{x'}(y') - G_x(y'))\| \cdot \|y - y'\| \\ &\leq \sqrt{\sum_{i=1}^n \frac{\beta_i^2 \gamma_i^2}{\alpha_i^2}} \cdot \|x - x'\| \cdot \|y - y'\|, \end{aligned}$$

where the last inequality follows from Lemma 1 and the standing Assumptions 1 and 2. Dividing through by $\|y - y'\|$ guarantees that $\text{Nash}(\cdot)$ is indeed a contraction with parameter $\sqrt{\sum_{i=1}^n \frac{\beta_i^2 \gamma_i^2}{\alpha_i^2}}$. The result follows immediately from the Banach fixed point theorem. $\blacksquare$

4.2 Repeated Gradient Method

Repeated retraining is largely a conceptual algorithm since in each iteration it requires computation of the exact Nash equilibrium of a stochastic game (11). A more realistic

algorithm would instead take a single gradient step on the game (11). With this in mind, given a step-size parameter $\eta > 0$, the *repeated gradient method* repeats the updates:

$$x^{t+1} = \text{proj}_{\mathcal{X}}(x^t - \eta G_{x^t}(x^t)).$$

More explicitly, in iteration t , each player $i \in [n]$ performs the update

$$x_i^{t+1} = \text{proj}_{\mathcal{X}_i} \left(x^t - \eta_t \mathbb{E}_{z_i \sim \mathcal{D}_i(x^t)} \nabla_i \ell_i(x_i^t, x_{-i}^t, z_i) \right).$$

This algorithm is largely conceptual since each player needs to compute an expectation; nonetheless we next show that this process converges linearly under the following additional smoothness assumption.

Assumption 3 (Smoothness) *For every $y \in \mathcal{X}$, the vector of individual gradients $G_y(x)$ is L -Lipschitz continuous in x .*

The following is the main result of this section. Note that the slightly suboptimal parameter regime $\rho < \frac{1}{2+\sqrt{2}}$ can be widened to the regime $\rho < 1$ by a more involved argument (cf. Theorem 10 with $\sigma^2 = 0$).

Theorem 9 (Repeated gradient method) *Suppose that Assumptions 1-3 hold and that we are in the regime $\rho < \frac{1}{2+\sqrt{2}}$. Then the iterates x^t generated by the repeated gradient method with $\eta = \frac{\alpha}{L^2}$ converge linearly to the performatively stable equilibrium x^* —that is, the following estimate holds:*

$$\|x^{t+1} - x^*\| \leq \left(\frac{1}{\sqrt{1 + \frac{\alpha^2}{L^2}}} + \frac{\alpha^2 \rho}{L^2} \right) \|x^t - x^*\| \quad \text{for all } t \geq 0. \quad (13)$$

Proof Using the triangle inequality, we estimate

$$\begin{aligned} \|x^{t+1} - x^*\| &= \|\text{proj}_{\mathcal{X}}(x^t - \eta G_{x^t}(x^t)) - x^*\| \\ &\leq \|\text{proj}_{\mathcal{X}}(x^t - \eta G_{x^*}(x^t)) - x^*\| \\ &\quad + \|\text{proj}_{\mathcal{X}}(x^t - \eta G_{x^t}(x^t)) - \text{proj}_{\mathcal{X}}(x^t - \eta G_{x^*}(x^t))\| \\ &\leq \|\text{proj}_{\mathcal{X}}(x^t - \eta G_{x^*}(x^t)) - x^*\| + \eta \|G_{x^t}(x^t) - G_{x^*}(x^t)\| \\ &\leq \|\text{proj}_{\mathcal{X}}(x^t - \eta G_{x^*}(x^t)) - x^*\| + \eta \sqrt{\sum_i \beta_i^2 \gamma_i^2} \cdot \|x^t - x^*\|, \end{aligned} \quad (14)$$

where the last inequality follows from Lemma 5. The rest of the argument is standard. We will simply show that the first term on the on right-side is a fraction of $\|x^t - x^*\|$. To this end, set $y^{t+1} = \text{proj}_{\mathcal{X}}(x^t - \eta G_{x^*}(x^t))$. Since the function $x \mapsto \frac{1}{2} \|x^t - \eta G_{x^*}(x^t) - x\|^2$ is 1-strongly convex and y^{t+1} is its minimizer over $\mathcal{X}$, we deduce

$$\frac{1}{2} \|y^{t+1} - x^*\|^2 \leq \frac{1}{2} \|x^t - \eta G_{x^*}(x^t) - x^*\|^2 - \frac{1}{2} \|x^t - \eta G_{x^*}(x^t) - y^{t+1}\|^2.$$

Expanding the right hand side yields

$$\begin{aligned}
 \frac{1}{2}\|y_{t+1} - x^\star\|^2 &\leq \frac{1}{2}\|x^t - x^\star\|^2 - \eta\langle G_{x^\star}(x^t), y^{t+1} - x^\star \rangle - \frac{1}{2}\|y^{t+1} - x^t\|^2 \\
 &= \frac{1}{2}\|x^t - x^\star\|^2 - \eta\langle G_{x^\star}(y^{t+1}), y^{t+1} - x^\star \rangle \\
 &\quad - \eta\langle G_{x^\star}(x^t) - G_{x^\star}(y^{t+1}), y^{t+1} - x^\star \rangle - \frac{1}{2}\|y^{t+1} - x^t\|^2.
 \end{aligned} \tag{15}$$

Strong convexity of the loss functions ensures

$$\eta\langle G_{x^\star}(y^{t+1}), y^{t+1} - x^\star \rangle \geq \eta\langle G_{x^\star}(y^{t+1}) - G_{x^\star}(x^\star), y^{t+1} - x^\star \rangle \geq \alpha\eta\|y^{t+1} - x^\star\|^2. \tag{16}$$

Young's inequality in turn implies

$$\begin{aligned}
 \eta|\langle G_{x^\star}(x^t) - G_{x^\star}(y^{t+1}), y^{t+1} - x^\star \rangle| &\leq \frac{\|G_{x^\star}(x^t) - G_{x^\star}(y^{t+1})\|^2}{2L^2} + \frac{\eta^2 L^2 \|y^{t+1} - x^\star\|^2}{2} \\
 &\leq \frac{\|x^t - y^{t+1}\|^2}{2} + \frac{\eta^2 L^2 \|y^{t+1} - x^\star\|^2}{2},
 \end{aligned} \tag{17}$$

where the last inequality follows from Assumption 3. Putting the estimates (15)-(17) together yields

$$\frac{1}{2}\|y^{t+1} - x^\star\|^2 \leq \frac{1}{2}\|x^t - x^\star\|^2 - \frac{2\alpha\eta - \eta^2 L^2}{2}\|y^{t+1} - x^\star\|^2.$$

Rearranging gives $\|y^{t+1} - x^\star\|^2 \leq \frac{1}{1+2\alpha\eta - \eta^2 L^2}\|x^t - x^\star\|^2$. Returning to (14) we therefore conclude

$$\|x^{t+1} - x^\star\| \leq \left(\frac{1}{\sqrt{1 + 2\alpha\eta - \eta^2 L^2}} + \eta \sqrt{\sum_i \beta_i^2 \gamma_i^2} \right) \|x^t - x^\star\|.$$

Plugging in $\eta = \frac{\alpha}{L^2}$ yields the claimed estimate (13). An elementary argument shows that in the assumed regime $\rho < \frac{1}{2+\sqrt{2}}$, the term in the parentheses in (13) is indeed smaller than one. $\blacksquare$

4.3 Repeated Stochastic Gradient Method

As observed earlier, the repeated gradient method is still largely a conceptual algorithm since an expectation has to be computed in every iteration. We next analyze an implementable algorithm that approximates the expectation in each step of gradient with an unbiased estimator. Namely, in each iteration t of the *repeated stochastic gradient method*, each player $i \in [n]$ performs the following update:

$$\left\{ \begin{array}{l} \text{Sample } z_i^t \sim \mathcal{D}_i(x^t) \\ \text{Set } x_i^{t+1} = \text{proj}_{\mathcal{X}_i}(x_i^t - \eta \nabla_{\ell_i}(x_i^t, x_{-i}^t, z_i^t)) \end{array} \right\}. \tag{18}$$

We will analyze the method under the following standard variance assumption. Recall the notation $\mathcal{D}_\pi(y) := \mathcal{D}_1(y) \times \dots \times \mathcal{D}_n(y)$.

Assumption 4 (Finite variance) *There exists a constant $\sigma \geq 0$ satisfying*

$$\mathbb{E}_{z \sim \mathcal{D}_\pi(x)} \|G_x(x) - g(x, z)\|^2 \leq \sigma^2 \quad \text{for all } x \in \mathcal{X}.$$

Convergence analysis for the repeated stochastic gradient method follows from the following simple observation. Taking into consideration the equality $G_x(x) = \mathbb{E}_{z \sim \mathcal{D}_\pi(x)} g(x, z)$, Lemma 5 directly implies that

$$\|G_x(x) - G_{x^*}(x)\| \leq \alpha \rho \|x - x^*\|.$$

That is, we may interpret the repeated stochastic gradient method as a standard stochastic gradient algorithm applied to the static problem $\mathcal{G}(x^*)$ with a bias that is linearly bounded by the distance to x^* . With this realization, we can simply analyze the stochastic gradient method on a static problem—i.e., where the decision distribution is fixed—with this special form of bias. Appendix A does exactly that. In particular, by taking $G(x) = G_{x^*}(x)$, the following is a direct consequence of Theorem 24 in Appendix A. This is in fact one of the interesting results of the paper—i.e., that analyzing multiplayer performative prediction problems with different updates reduces to analyzing stochastic gradient methods with different types of bias.

In the following theorem, let $\mathbb{E}_t$ denote the conditional expectation with respect to the σ -algebra generated by $(x^l)_{l=1, \dots, t}$.

Theorem 10 (One-step improvement) *Suppose that Assumptions 1-4 hold and that we are in the regime $\rho < 1$. Then with any $\eta < \frac{\alpha(1-\rho)}{8L^2}$, the repeated stochastic gradient method generates a sequence x^t satisfying*

$$\mathbb{E}_t \|x^{t+1} - x^*\|^2 \leq \frac{1 + 2\eta\alpha\rho + 2\eta^2\alpha^2\rho^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})} \|x^t - x^*\|^2 + \frac{4\eta^2\sigma^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})},$$

where x^* is the performatively stable equilibrium of the game (4).

Proof This follows directly from applying Theorem 24 in Appendix A with $G(x) = G_{x^*}(x)$, $g^t = g(x^t, z^t)$, $C_t = D = 0$, and $B = \alpha\rho$. $\blacksquare$

With Theorem 10 at hand, it is straightforward to obtain efficiency estimates—i.e. independent of the initialization—under a variety of step-size choices. One particular choice, highlighted by (Ghadimi and Lan, 2013), is the step-decay schedule that periodically cuts η by a fraction. This choice of the step-size schedule allows us to separate the rate of convergence into a deterministic part (as if there is zero variance) and the stochastic part. The deterministic part exhibits a standard linear rate of convergence where as the stochastic part exhibits a sublinear rate. Other step-size schedules do not typically lead to such a decomposition. The resulting algorithm and its convergence guarantees are summarized in the following corollary.

Corollary 11 (Repeated stochastic gradient method with a step-decay schedule) *Suppose that Assumptions 1-4 hold and we are in the regime $\rho < 1$. Set $\eta_0 := \frac{\alpha(1-\rho)}{4}$.*

$\min\{1, \frac{1}{2L^2}\}$. Consider running the repeated stochastic gradient method in $k = 0, \dots, K$ epochs, for T_k iterations each, with constant step-size $\eta_k = 2^{-k}\eta_0$, and such that the last iterate of epoch k is used as the first iterate in epoch $k + 1$. Fix a target accuracy $\varepsilon > 0$ and suppose we have available a constant $R \geq \|x^0 - x^*\|^2$. Set

$$T_0 = \left\lceil \frac{10}{(1-\rho)\alpha\eta_0} \log\left(\frac{2R}{\varepsilon}\right) \right\rceil, \quad T_k = \left\lceil \frac{10 \log(4)}{(1-\rho)\alpha\eta_k} \right\rceil \quad \text{for } k \geq 1, \quad \text{and} \quad K = \left\lceil 1 + \log_2\left(\frac{40\eta_0\sigma^2}{(1-\rho)\alpha\varepsilon}\right) \right\rceil.$$

The final iterate x produced satisfies $\mathbb{E} \|x - x^*\|^2 \leq \varepsilon$, while the total number of iterations of the repeated stochastic gradient method is at most

$$\mathcal{O}\left(\frac{L^2}{(1-\rho)\alpha^2} \cdot \log\left(\frac{2R}{\varepsilon}\right) + \frac{\sigma^2}{(1-\rho)^2\alpha^2\varepsilon}\right).$$

Proof Consider a sequence $x^0, x^1, \dots, x^t$ generated by the stochastic gradient method with a fixed step-size $\eta \leq \eta_0$. Using Theorem 10 together with the tower rule for conditional expectations, we deduce

$$\mathbb{E} \|x^{t+1} - x^*\|^2 \leq \frac{1 + 2\eta\alpha\rho + 2\eta^2\alpha^2\rho^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})} \mathbb{E} \|x^t - x^*\|^2 + \frac{4\eta^2\sigma^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})}. \quad (19)$$

Our choice of η_0 ensures

$$\frac{1 + 2\eta\alpha\rho + 2\eta^2\alpha^2\rho^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})} \leq \frac{1 + 2\eta\alpha \cdot \frac{1+3\rho}{4}}{1 + 2\eta\alpha \cdot \frac{1+\rho}{2}} = 1 - \frac{2\eta\alpha(\frac{1-\rho}{4})}{1 + 2\eta\alpha(\frac{1+\rho}{2})} \leq 1 - \frac{1-\rho}{10}\eta\alpha.$$

Therefore iterating (19) we obtain the estimate

$$\mathbb{E} \|x_{t+1} - x^*\|^2 \leq (1 - \psi(\eta))^{t+1} \|x_0 - x^*\|^2 + \Gamma\eta,$$

where we set $\psi(\eta) = c\eta$ with $c = \frac{1-\rho}{10}\alpha$ and $\Gamma = \frac{40\sigma^2}{\alpha(1-\rho)}$. The result now follows directly from (Drusvyatskiy and Xiao, 2023, Lemma B.2). $\blacksquare$

The efficiency estimate in Corollary 11 coincides with the standard efficiency estimate for the stochastic gradient method on static problems, up to multiplication by $(1 - \rho)^{-2}$.

Remark 12 (Asynchronous Feedback) In practice, it may not be the case that the decision makers observe data or actions synchronously, and as a result they may not have the requisite information to update their action in every time step (iteration). A natural model to capture asynchronous updates is one in which decision maker i receives sufficient information to update its decision x_i with probability p_i . For instance, in the case of the repeated stochastic gradient method, this means that

$$x_i^{t+1} = \begin{cases} \text{proj}_{\mathcal{X}_i}(x_i^t - \eta \nabla_i \ell_i(x_i^t, x_{-i}^t, z_i^t)), & \text{w.p. } p_i \\ x_i^t, & \text{w.p. } (1 - p_i) \end{cases} \quad (20)$$

This type of update has been studied fairly extensively in the literature on stochastic optimization and in learning in games—see, e.g., (Recht et al., 2011; Lian et al., 2015;

Huo and Huang, 2017; Mertikopoulos and Zhou, 2019; Zhou et al., 2018; Chasnov et al., 2020b) and references therein. In the context of finding performatively stable points via the repeated stochastic gradient method in (18) modified via (20), the rates in Corollary 11 do not change much; the primary difference is that the Lipschitz constants are rescaled by $p_{\max} := \max\{p_1, \dots, p_n\}$ and the strong monotonicity constant is rescaled by $p_{\min} := \min\{p_1, \dots, p_n\}$. The reason this works out is that we can simply perform the exact same analysis using a modified inner product as has been performed in prior literature—i.e., we simply perform the analysis in the inner product $[x, y] = \langle P^{-1}x, y \rangle$ where $P = \text{diag}(p_1, \dots, p_n)$. Since there are new insights deriving from the performative structure of the problem in the convergence analysis beyond what we have already shown for the synchronous case, for brevity we do not go into the full details for the asynchronous case; the results follow precisely from the analysis we have already presented except in this new coordinate system.

5. Monotonicity of Decision-Dependent Games

Our next goal is to develop algorithms for finding true Nash equilibria of the game (4). As the first step, this section presents sufficient conditions for the game to be monotone along with some examples. We note, however, that the sufficient conditions we present are strong because the game (4) is typically not monotone. When specialized to the single player setting $n = 1$, the sufficient conditions we derive are identical to those in (Miller et al., 2021) although the proofs are entirely different.

We impose the following mild smoothness assumption.

Assumption 5 (Smoothness of the distribution) *For each index $i \in [n]$ and point $x \in \mathcal{X}$, the map $u_i \mapsto \mathbb{E}_{z_i \sim \mathcal{D}(u_i, x_{-i})} \ell_i(x, z_i)$ is differentiable at $u_i = x_i$ and its derivative is continuous in x .*

Under Assumption 5, direct differentiation implies the following expression for the derivative of player i 's loss function which is precisely (6) from Section 3:

$$\nabla_i \mathcal{L}_i(x_i, x_{-i}) = \frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(x_i, x_{-i})} [\ell_i(u_i, x_{-i}, z_i)] \Big|_{u_i=x_i} + \frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x_{-i})} [\ell_i(x_i, x_{-i}, z_i)] \Big|_{u_i=x_i}.$$

To simplify notation, define

$$H_{i,x}(y) := \frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}(u_i, x_{-i})} \ell_i(y, z_i) \Big|_{u_i=x_i}.$$

Therefore, we may equivalently write

$$\nabla_i \mathcal{L}_i(x_i, x_{-i}) = G_{i,x}(x) + H_{i,x}(x)$$

where $G_{i,x}(x)$ is defined in (7). Stacking together the individual partial gradients $H_{i,x}(y)$ for each player, we set $H_x(y) = (H_{1,x}(y), \dots, H_{n,x}(y))$. Therefore the vector of individual gradients for (4) is simply the map $D(x) := G_x(x) + H_x(x)$. Thus the game (4) is monotone, as long as $D(x)$ is a monotone mapping.

The sufficient conditions for monotonicity (cf. Theorem 13 below) are simply that we are in the regime $\rho < \frac{1}{2}$ and that the map $x \mapsto H_x(y)$ is monotone for any y . The latter can

be understood as requiring that for any $y \in \mathcal{X}$, the auxiliary game wherein each player aims to solve

$$\min_{x_i \in \mathcal{X}} \mathbb{E}_{z_i \sim \mathcal{D}_i(x_i, x_{-i})} \ell_i(y, z_i).$$

is monotone. In the single player setting $n = 1$, this simply means that the function $x \mapsto \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} \ell(y, z_i)$ is convex for any fixed $y \in \mathcal{X}$, thereby reducing exactly to the requirement in (Miller et al., 2021, Theorem 3.1).

Theorem 13 (Monotonicity of the decision-dependent game) *Suppose that Assumptions 1, 2, and 5 hold, and that we are in the regime $\rho < \frac{1}{2}$ and the map $x \mapsto H_x(y)$ is monotone for any y . The game (4) is strongly monotone with parameter $(1 - 2\rho)\alpha$.*

The proof of Theorem 13 crucially relies on the following useful lemma.

Lemma 14 *Suppose that Assumptions 1, 2, and 5 hold. For any $x \in \mathcal{X}$, the map $H_x(y)$ is Lipschitz continuous in y with parameter $\sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}$.*

Proof Fix three points $x, x', y \in \mathcal{X}$. Player i 's coordinate of $H_{x'}(x) - H_{x'}(y)$ is simply

$$H_{i,x'}(x) - H_{i,x'}(y) = \frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x'_{-i})} (\ell_i(x, z_i) - \ell_i(y, z_i)) \Big|_{u_i=x'_i}.$$

Setting $\gamma(s) = y + s(x - y)$ for any $s \in (0, 1)$, the fundamental theorem of calculus ensures

$$\ell_i(x, z_i) - \ell_i(y, z_i) = \int_{s=0}^1 \langle \nabla_i \ell_i(\gamma(s), z_i), x - y \rangle ds.$$

Therefore differentiating, taking an expectation, and using the Cauchy-Schwarz inequality we deduce

$$\|H_{i,x'}(x) - H_{i,x'}(y)\| \leq \int_{s=0}^1 \left\| \frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x'_{-i})} \nabla_i \ell_i(\gamma(s), z_i) \Big|_{u_i=x'_i} \right\| \cdot \|x - y\| ds. \quad (21)$$

Now for any $s \in (0, 1)$, Lemma 5 guarantees that the map $u_i \mapsto \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x'_{-i})} \nabla_i \ell_i(\gamma(s), z_i)$ is Lipschitz continuous with parameter $\beta_i \gamma_i$ and therefore its derivative is upper-bounded in norm by $\beta_i \gamma_i$. We therefore deduce that the right hand side of (21) is upper bounded by $\beta_i \gamma_i \|x - y\|$. Applying this argument to each player leads to the claimed Lipschitz constant on $H_x(y)$ with respect to x . $\blacksquare$

Given the preceding lemma, we now prove Theorem 13.

Proof [Proof of Theorem 13] Fix an arbitrary pair of points $x, x' \in \mathcal{X}$. Expanding the following inner product, we have

$$\langle D(x) - D(x'), x - x' \rangle = \langle G_x(x) - G_{x'}(x'), x - x' \rangle + \langle H_x(x) - H_{x'}(x'), x - x' \rangle. \quad (22)$$

We estimate the first term as follows:

$$\begin{aligned} \langle G_x(x) - G_{x'}(x'), x - x' \rangle &= \langle G_{x'}(x) - G_{x'}(x'), x - x' \rangle + \langle G_x(x) - G_{x'}(x), x - x' \rangle \\ &\geq \alpha \|x - x'\|^2 - \left(\sum_{i=1}^n \beta_i^2 \gamma_i^2 \right)^{1/2} \cdot \|x - x'\|^2 \end{aligned} \quad (23)$$

$$= (1 - \rho) \alpha \|x - x'\|^2, \quad (24)$$

where (23) follows from Assumption 1 and Lemma 5. Next, we estimate the second term on the right side of (22) as follows:

$$\begin{aligned} \langle H_x(x) - H_{x'}(x'), x - x' \rangle &= \langle H_{x'}(x) - H_{x'}(x'), x - x' \rangle + \langle H_x(x) - H_{x'}(x), x - x' \rangle \\ &\geq \langle H_{x'}(x) - H_{x'}(x'), x - x' \rangle \end{aligned} \quad (25)$$

$$\geq -\|H_{x'}(x) - H_{x'}(x')\| \cdot \|x - x'\| \quad (26)$$

$$\geq -\left(\sum_{i=1}^n \beta_i^2 \gamma_i^2 \right)^{1/2} \|x - x'\|^2, \quad (27)$$

where (25) follows from the assumption that the map $x \mapsto H_x(y)$ is monotone and (27) follows from Lemma 14. Combining (22), (24), and (27) completes the proof. $\blacksquare$

Observe that for the game to be strongly monotone we need that the map $D(x) = G_x(x) + H_x(x)$ is strongly monotone. In Section 4, to obtain convergence results we simply need that $G_x(x)$ is strongly monotone, and hence, under the assumption that $G_x(x)$ is strongly monotone, in order for $D(x)$ to be strongly monotone, it is sufficient for $H_x(x)$ to be monotone. That being said, even if $G_x(x)$ is not strongly monotone, as long as $H_x(x)$ is sufficiently monotone, then $D(x)$ will be. Thus the conditions we provide in this section are merely sufficient. Below we provide examples of performative prediction games that may arise in applications of machine learning, and provide conditions under which strongly monotonicity holds.

Indeed, the following two examples of multiplayer performative prediction problems illustrate settings where the mapping $x \mapsto H_x(y)$ is indeed monotone and therefore Theorem 13 may be applied to deduce monotonicity of the game. We explore both these examples numerically in Section 7.

Example 1 (Revenue Maximization in Rideshare Markets) Consider a rideshare market with two firms that each would like to maximize their revenue by setting the price x_i . The demand z_i that each ride share firm sees is influenced not only by the price they set but also the price that their competitor sets. Suppose that firm i 's loss is given by

$$\ell_i(x_i, z_i) = -z_i^\top x_i + \frac{\lambda_i}{2} \|x_i\|^2$$

where $\lambda_i \geq 0$ is some regularization parameter. Moreover, let us suppose that the random demand z_i takes the semi-parametric form $z_i = \zeta_i + A_i x_i + A_{-i} x_{-i}$, where ζ_i follows some base distribution $\mathcal{P}_i$ and the parameters A_i and A_{-i} capture price elasticities to player i 's

and its competitor's change in price, respectively. The mapping $x \mapsto H_x(y)$ is monotone. Indeed, observe that i -th component of $H_x(y)$ is given by

$$H_{i,x}(y) = \mathbb{E}_{\zeta_i \sim \mathcal{P}_i} A_i^\top \nabla_{z_i} \ell_i(y_i, \zeta_i + A_i x_i + A_{-i} x_{-i}) = -A_i^\top y_i.$$

Hence, the map $x \mapsto H_x(y)$ is constant and is therefore trivially monotone.

The next example is a multiplayer extension of Example 3.2 in (Miller et al., 2021) which models the single player decision-dependent problem of predicting the final vote margin in an election contest.

Example 2 (Strategic Prediction) Consider two election prediction platforms. Each platform seeks to predict the vote margin. Not only can predicting a large margin in either direction dissuade people from voting, but people may look at multiple platforms as a source for information. Features θ such as past polling averages are drawn i.i.d. from a static distribution $\mathcal{P}_\theta$. Each platform observes a sample drawn from the conditional distribution

$$z_i | \theta \sim \varphi_i(\theta) + A_i x_i + A_{-i} x_{-i} + w_i,$$

where $\varphi_i : \mathbb{R}^{d_i} \rightarrow \mathbb{R}$ is an arbitrary map, the parameter matrices $A_i \in \mathbb{R}^{m_i \times d_i}$ and $A_{-i} \in \mathbb{R}^{m_i \times d_{-i}}$ are fixed, and w_i is a zero-mean random variable. Each player seeks to predict z_i by minimizing the loss

$$\ell_i(x_i, z_i) = \frac{1}{2} \|z_i - \theta^\top x_i\|^2.$$

We claim that the map $x \mapsto H_x(y)$ is monotone as long as

$$\sqrt{n-1} \cdot \max_{i \in [n]} \|A_{-i}^\top A_i\|_{\text{op}} \leq \min_{i \in [n]} \lambda_{\min}(A_i^\top A_i),$$

where $\lambda_{\min}$ denotes the minimal eigenvalue. The interpretation of this condition is that the performative effects due to interaction with competitors are small relative to any player's own performative effects. To see the claim, set $\bar{A}_i$ to be the matrix satisfying $\bar{A}_i x = A_i x_i + A_{-i} x_{-i}$ and observe that the i -th component of $H_x(y)$ is given by

$$\begin{aligned} H_{i,x}(y) &= \mathbb{E}_{\theta, w_i} A_i^\top \nabla_{z_i} \ell_i(y_i, \varphi_i(\theta) + \bar{A}_i x + w_i) \\ &= \mathbb{E}_{\theta, w_i} A_i^\top (\varphi_i(\theta) + \bar{A}_i x - \theta^\top y_i + w_i) \\ &= A_i^\top A_i x_i + A_i^\top A_{-i} x_{-i} + \mathbb{E}_{\theta, w_i} A_i^\top (\varphi_i(\theta) - \theta^\top y_i + w_i). \end{aligned}$$

Therefore, the map $H_{i,x}(y)$ is affine in x . Consequently, monotonicity of $x \mapsto H_{i,x}(y)$ is equivalent to monotonicity of the linear map $x \mapsto V(x) + W(x)$, where V is the block diagonal matrix $V(x) = \text{Diag}(A_1^\top A_1, \dots, A_n^\top A_n)x$ and we define the linear map $W(x) = (A_1^\top A_{-1} x_{-1}, \dots, A_n^\top A_{-n} x_{-n})$. The minimal eigenvalue of V is simply $\min_{i \in [n]} \lambda_{\min}(A_i^\top A_i)$. Let us estimate the operator norm of W . To this end, set $s := \max_i \|A_i^\top A_{-i}\|_{\text{op}}$ and for any x we compute

$$\|W(x)\|^2 = \sum_{i=1}^n \|A_i^\top A_{-i} x_{-i}\|^2 \leq \sum_{i=1}^n s^2 \|x_{-i}\|^2 = (n-1)s^2 \|x\|^2.$$

Thus, under the stated assumptions, the operator norm of W is smaller than the minimal eigenvalue of V , and therefore the sum $V + W$ is monotone.

6. Algorithms for Finding Nash Equilibria

In contrast to Section 4, in this section we analyze algorithms that converge to the Nash equilibrium of the n -player performative prediction game (4) when the game is strongly monotone. Recall that the Nash equilibrium x^* of this game is characterized by the relation

$$x_i^* \in \operatorname{argmin}_{x_i \in \mathcal{X}_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(x_i, x_{-i}^*)} \ell_i(x_i, x_{-i}^*, z_i) \quad \forall i \in [n].$$

It is important to stress the distinction between the performatively stable equilibria studied in Section 4 and the Nash equilibrium x^* of the game (4): namely, for the latter concept the distribution $\mathcal{D}_i$ explicitly depends on the optimization variable x_i versus being fixed at $x^* = (x_i^*, x_{-i}^*)$. Theorem 13 (cf. Section 5) gives sufficient conditions under which the multiplayer performative prediction game (4) is strongly monotone and hence, admits a unique Nash equilibrium.

In the following subsections, we study natural learning dynamics—namely, variants of *gradient play* as it is referred to in the literature on learning and games—for continuous games seeking Nash equilibrium in different information settings. Specifically, we study gradient-based learning methods where players update using an estimate of their individual gradient consistent with the information available to them. It is important to contrast the gradient updates in Section 4 with the updates considered in this section: the Nash-seeking algorithms studied in this section all use gradient estimates of the individual gradient

$$\nabla_i \mathcal{L}_i(x_i, x_{-i}) = \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} [\nabla_i \ell_i(x, z_i)] + \frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x_{-i})} [\ell_i(x, z_i)] \Big|_{u_i=x_i} \quad (28)$$

for each player $i \in [n]$, whereas performatively stable equilibrium seeking algorithms of Section 4 are defined such that the gradient update only uses the first term on the right hand side of (28).

The main difficulty with applying gradient-based methods is that estimation of the second term on the right hand side of (28), without some parametric assumptions on the distributions $\mathcal{D}_i$. Consequently, we start in Section 6.1 with derivative free methods, wherein each player only has access to loss function queries. This does not require players to have any information on the distribution $\mathcal{D}_i$, but results in a slow algorithm, roughly with complexity $\mathcal{O}(\frac{d^2}{\varepsilon^2})$. In practice, the players may have some information on $\mathcal{D}_i$, and it's reasonable that they would exploit this information during learning. Hence, in Section 6.2 we impose a specific parametric assumption of the distributions and study *stochastic gradient play*³ under the assumption that each player knows their own “influence” parameter on the distribution. The resulting algorithm enjoys efficiency on the order of $\mathcal{O}(\frac{1}{\varepsilon})$. Section 6.3 instead develops a variant of a stochastic gradient method wherein each player adaptively learns their influence parameters, and uses their current estimate of those parameters to optimizing their loss function by taking a step along the direction of their individual gradient; the resulting process has efficiency on the order of $\mathcal{O}(\frac{d}{\varepsilon})$.

3. This method is known as stochastic gradient play in the game theory literature; here we refer to as the *stochastic gradient method* to be consistent with the naming convention of other methods in the paper.

Remark 15 (Asynchronous Feedback Revisted) As noted in Remark 12, in practice, it may not be the case that the decision makers observe data or actions synchronously. In the pursuit of learning Nash equilibrium in this setting, we can adopt the same model as described in Remark 12—wherein player i receives gradient information with probability p_i in each iteration—to capture asynchronous information feedback. As noted, this type of update has been studied fairly extensively in the literature on stochastic optimization and in learning in games. In the context of finding Nash equilibrium via the methods proposed in the remainder of this section, up to scaling factors proportional to $p_{\min} := \min\{p_1, \dots, p_n\}$ and $p_{\max} := \max\{p_1, \dots, p_n\}$, the rates we present do not fundamentally change. Again, this follows simply by using the exact same analysis in a transformed coordinate system, and so for brevity we do not go into the full details for the asynchronous case.

6.1 Derivative Free Method for Performative Prediction Games

As just alluded to, the first information setting we consider for multiplayer performative prediction is the “bandit feedback” setting, where players have oracle access to queries of their loss function only, and therefore are faced with the problem of creating an estimate of their gradient from such queries. This setting requires the least assumptions on what information is available to players. In the optimization literature, when a first order oracle is not available, derivative free or zeroth order methods are typically applied. Derivative free methods have been extended to games (Bravo et al., 2018; Drusvyatskiy et al., 2022). The results in this section are direct consequences of the results in these papers. We concisely spell them out here in order to compare them with the convergence guarantees discussed in the following two sections.

The *derivative free (gradient) method* we consider proceeds as follows. Fix a parameter $\delta > 0$. In each iteration t , each player $i \in [n]$ performs the following update:

$$\left\{ \begin{array}{l} \text{Sample } v_i^t \in \mathbb{S}_i \\ \text{Sample } z_i^t \sim \mathcal{D}_i(x^t + \delta v^t) \\ \text{Set } x_i^{t+1} = \text{proj}_{(1-\delta)\mathcal{X}_i} \left(x_i^t - \eta_t \frac{d_i}{\delta} \ell_i(x^t + \delta v^t, z_i^t) v_i^t \right) \end{array} \right\}. \quad (29)$$

There are not additional information requirements on the players beyond what is assumed for derivative gradient play as studied for example in (Bravo et al., 2018; Drusvyatskiy et al., 2022). Each player simply submits its query $x_i^t + \delta v_i^t$ to the environment and receives back its loss $\ell_i(x^t + \delta v^t, z_i^t)$ where $z_i^t \sim \mathcal{D}(x^t + \delta v^t)$.

Recall that $\mathbb{S}_i$ denotes the unit sphere with dimension d_i . The reason for projecting onto the set $(1 - \delta)\mathcal{X}_i$ is simply to ensure that in the next iteration $t + 1$, the strategy played by player i remains in $\mathcal{X}_i$. We state the convergence guarantees of the method informally here because they are meant only as a baseline result. The formal statement for derivative free methods in general games can be found in Drusvyatskiy et al. (2022).⁴

4. Though Theorem 2 in Drusvyatskiy et al. (2022) is stated for deterministic games, it applies verbatim whenever the value of the loss function for each each player is replaced by an unbiased estimator of their individual loss functions—our setting.

Proposition 16 (Convergence rate of the derivative free method) *Consider an n -player decision-dependent game as defined in (4). Under reasonable smoothness and bounded variance assumptions, algorithm (29) with appropriately chosen parameters δ and η_t will find a point x satisfying $\mathbb{E}[\|x - x^*\|^2] \leq \varepsilon$ after at most $O(\frac{d^2}{\varepsilon^2})$ iterations.*

The rate $O(\frac{d^2}{\varepsilon^2})$ can be extremely slow in practice. The constants in this rate may be improved slightly in our setting by taking into consideration that the gradient (cf. (28)) is composed of two terms, one of which is known and the other which is not: $\mathbb{E}_{z_i \sim \mathcal{D}_i(x)}[\nabla_i \ell_i(x, z_i)]$ which is known, and $\frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x_{-i})}[\ell_i(x, z_i)]|_{u_i=x_i}$ which is unknown without *a priori* knowledge on the structure of $\mathcal{D}_i(\cdot)$. Therefore, if player i only applies a derivative free estimation scheme on the second part of the gradient then the constants showing up in the rate could be improved.

In the rest of the paper, we focus on stochastic gradient based methods which enjoy significantly better efficiency guarantees (at cost of access to a richer oracle).

6.2 Stochastic Gradient Method in Performative Prediction Games

In practice, often players have some information regarding their data distribution $\mathcal{D}_i$ and can leverage this during learning. Stochastic gradient play—which we refer to as the stochastic gradient method to be consistent with the rest of the paper—is a natural learning algorithm commonly adopted in the literature on learning in games for settings where players have an unbiased estimate of their individual gradient. To apply the stochastic gradient method to multiplayer performative prediction, players need oracle access to the gradient of their loss with respect to their choice variable, which requires some knowledge of how the distribution $\mathcal{D}_i$ depends on the joint action profile x . To this end, let us impose the following parametric assumption, which we have already encountered in Example 1.

Assumption 6 (Parametric assumption) *For each index $i \in [n]$, there exists a probability measure $\mathcal{P}_i$ and matrices A_i and A_{-i} satisfying*

$$z_i \sim \mathcal{D}_i(x) \iff z_i = \zeta_i + A_i x_i + A_{-i} x_{-i} \quad \text{for } \zeta_i \sim \mathcal{P}_i.$$

The mean and covariance of ζ_i are defined as $\mu_i := \mathbb{E}_{\zeta_i \sim \mathcal{P}_i}[\zeta_i]$ and $\Sigma_i := \mathbb{E}_{\zeta_i \sim \mathcal{P}_i}[(\zeta_i - \mu_i)(\zeta_i - \mu_i)^\top]$, respectively.

Assumption 6 generalizes an analogous modeling very commonly used in the single player performative prediction setting in (Miller et al., 2021). It can also be viewed as a local linear approximation of nonlinear behavior. It asserts that the distribution used by player i is a “linear perturbation” of some base distribution $\mathcal{P}_i$. We can interpret the matrices A_i and A_{-i} as quantifying the performative effects of player i ’s decisions and the rest of the players’ decisions, respectively, on the distribution $\mathcal{D}_i$ governing player i ’s data.

Under Assumption 6, we may write player i -th loss function as

$$\mathcal{L}_i(x) = \mathbb{E}_{\zeta_i \sim \mathcal{P}_i} \ell_i(x, \zeta_i + A_i x_i + A_{-i} x_{-i}). \quad (30)$$

Under mild smoothness assumptions, differentiating (30) using the chain rule, we see that the gradient of the i -th player’s loss is simply

$$\nabla_i \mathcal{L}_i(x) = \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} [\nabla_i \ell_i(x, z_i) + A_i^\top \nabla_{z_i} \ell_i(x, z_i)]. \quad (31)$$

Therefore, given a point x , player i may draw $z_i \sim \mathcal{D}_i(x)$ and form the vector

$$w_i(x, z_i) = \nabla_i \ell_i(x, z_i) + A_i^\top \nabla_{z_i} \ell_i(x, z_i).$$

By definition, $w_i(x, z_i)$ is an unbiased estimator of $\nabla_i \mathcal{L}_i(x)$, that is

$$\mathbb{E}_{z_i \sim \mathcal{D}_i(x)} w_i(x, z_i) = \nabla_i \mathcal{L}_i(x).$$

With this notation, the stochastic gradient method proceeds as follows: in each iteration $t \geq 0$ each player $i \in [n]$ performs the update:

$$\left\{ \begin{array}{l} \text{Sample } z_i^t \sim \mathcal{D}_i(x^t) \\ \text{Set } x_i^{t+1} = \text{proj}_{\mathcal{X}_i}(x_i^t - \eta_t \cdot w_i(x^t, z_i^t)) \end{array} \right\}. \quad (32)$$

Let us look at the computation that is required in each iteration. Evaluation of the vector $w_i(x, z_i)$ requires evaluation of both $\nabla_i \ell_i(x, z_i)$ and $\nabla_{z_i} \ell_i(x, z_i)$, and knowledge of the matrix A_i . When the game is separable, it is very reasonable that each player can explicitly compute $\nabla_i \ell_i(x_i, z_i)$ and $\nabla_{z_i} \ell_i(x_i, z_i)$ assuming oracle access to queries z_i from the environment which does depend on x_{-i} and x_i . Moreover, the matrix A_i depends only on the performative effects of player i , and in this section we will suppose that it is indeed known to each player. In the next section, we will develop an adaptive algorithm wherein each player $i \in [n]$ simultaneously learns A_i and A_{-i} while optimizing their loss.

In order to apply standard convergence guarantees for stochastic gradient play, we need to assume that (i) the vector of individual gradients is Lipschitz continuous and (ii) that the variance of $w(x, z_i)$ is bounded. Let us begin with the former.

Assumption 7 (Smoothness) *Suppose that the map $(\nabla_1 \mathcal{L}_1(x), \nabla_2 \mathcal{L}_2(x), \dots, \nabla_n \mathcal{L}_n(x))$ is L -Lipschitz continuous.*

The constant L may be easily estimated from the smoothness parameters of each individual loss function $\ell_i(x, z)$ and the magnitude of the matrices A_i and A_{-i} ; this is the content of the following lemma. In what follows, we define the mixed partial derivative $\nabla_{i,z_i} \ell_i(x, z_i) = (\nabla_i \ell_i(x, z_i), \nabla_{z_i} \ell_i(x, z_i))$. Recall that $\nabla_i \ell_i(x_i, x_{-i}, z_i)$ denotes the partial derivative of ℓ_i with respect to the x_i argument and $\nabla_{z_i} \ell_i(x_i, x_{-i}, z_i)$ denotes the partial derivative with respect to z_i .

Lemma 17 (Sufficient conditions for Assumption 7) *Suppose that Assumption 6 holds and that there exist constants $\xi_i \geq 0$ such that for each index i the map $(x, z_i) \mapsto \nabla_{i,z_i} \ell_i(x, z_i)$ is ξ_i -Lipschitz continuous. Then Assumption 7 holds with*

$$L = \sqrt{\sum_{i=1}^n \xi_i^2 \max\{1, \|A_i\|_{\text{op}}^2\} \cdot (1 + \|\bar{A}_i\|_{\text{op}}^2)}.$$

Proof Let $\bar{A}_i$ be a matrix satisfying $\bar{A}_i x = A_i x_i + A_{-i} x_{-i}$. Observe that we may write

$$\nabla_i \mathcal{L}_i(x) = \mathbb{E}_{\zeta_{i,0} \sim \mathcal{P}_i} V^\top \nabla_{i,z_i} \ell_i(x, \zeta_i + \bar{A}_i x) \quad \text{where} \quad V = \begin{bmatrix} I & 0 \\ 0 & A_i \end{bmatrix}.$$

Therefore, we deduce

$$\begin{aligned}
 \|\nabla_i \mathcal{L}_i(x) - \nabla_i \mathcal{L}_i(x')\| &\leq \|V\|_{\text{op}} \mathbb{E}_{\zeta_i \sim \mathcal{P}_i} \|\nabla_{i, \zeta_i} \ell_i(x, \zeta_i + \bar{A}_i x) - \nabla_{i, \zeta_i} \ell_i(x', \zeta_i + \bar{A}_i x')\| \\
 &\leq \max\{1, \|A_i\|_{\text{op}}\} \cdot \xi_i \cdot \mathbb{E}_{\zeta_i \sim \mathcal{P}_i} \|(x, \zeta_i + \bar{A}_i x) - (x', \zeta_i + \bar{A}_i x')\| \\
 &= \max\{1, \|A_i\|_{\text{op}}\} \cdot \xi_i \cdot \sqrt{\|x - x'\|^2 + \|\bar{A}_i(x - x')\|^2} \\
 &\leq \max\{1, \|A_i\|_{\text{op}}\} \cdot \xi_i \cdot \sqrt{1 + \|\bar{A}_i\|_{\text{op}}^2} \cdot \|x - x'\|.
 \end{aligned}$$

This completes the proof. $\blacksquare$

Next, we assume a finite variance bound.

Assumption 8 (Finite variance) *Suppose that there exists a constant $\sigma > 0$ satisfying*

$$\mathbb{E}_{z \sim \mathcal{D}_\pi(x)} \|w(x, z) - \mathbb{E}_{z' \sim \mathcal{D}_\pi(x)} w(x, z')\|^2 \leq \sigma^2 \quad \forall x \in \mathcal{X}.$$

Let us again present a sufficient condition for Assumption 8 to hold in terms of the variance of the individual gradients $\nabla_{i, z_i} \ell_i(x, z_i)$. The proof is immediate and we omit it.

Lemma 18 (Sufficient conditions for Assumption 8) *Suppose that there exist constants $s_1, s_2 \geq 0$ such that for all $x \in \mathcal{X}$ and $i \in [n]$ the estimates hold:*

$$\begin{aligned}
 \mathbb{E}_{z'_i \sim \mathcal{D}_i(x)} \|\nabla_i \ell_i(x, z'_i) - \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} \nabla_i \ell_i(x, z_i)\|^2 &\leq s_1^2 \\
 \mathbb{E}_{z'_i \sim \mathcal{D}_i(x)} \|\nabla_{z_i} \ell_i(x, z'_i) - \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} \nabla_{z_i} \ell_i(x, z_i)\|^2 &\leq s_2^2
 \end{aligned}$$

Then Assumption 8 holds with $\sigma^2 = \sum_{i=1}^n 2(s_1^2 + \|A_i\|_{\text{op}}^2 s_2^2)$.

The following is now a direct consequence of standard convergence guarantees for stochastic gradient methods.

Theorem 19 (Stochastic gradient play) *Consider an n -player performative prediction game (4). Suppose that Assumptions 6-8 hold and that the game is α -strongly monotone with $\alpha > 0$. Then a single step of the stochastic gradient method (32) with any constant $\eta \leq \frac{\alpha}{2L^2}$ satisfies*

$$\mathbb{E}[\|x^{t+1} - x^*\|^2] \leq \frac{1}{1 + \alpha\eta} \mathbb{E}[\|x^t - x^*\|^2] + \frac{2\eta^2\sigma^2}{1 + \eta\alpha}, \quad (33)$$

where x^* is the Nash equilibrium of the game (4).

Proof This is immediate from Theorem 24 in Appendix A with $B \equiv C_t \equiv D \equiv 0$. $\blacksquare$

Analogous to the analysis of the stochastic repeated gradient method, applying a step-decay schedule on η yields the following corollary. The proof follows directly from the recursion (33) and the generic results on step decay schedules; e.g. (Drusvyatskiy and Xiao, 2023, Lemma B.2).

Corollary 20 (Stochastic gradient method with a step-decay schedule) *Suppose that the assumptions of Theorem 19 hold. Consider running stochastic gradient method in $k = 0, \dots, K$ epochs, for T_k iterations each, with constant step-size $\eta_k = \frac{\alpha}{2L^2} \cdot 2^{-k}$, and such that the last iterate of epoch k is used as the first iterate in epoch $k+1$. Fix a target accuracy $\varepsilon > 0$ and suppose we have available a constant $R \geq \|x^0 - x^*\|^2$. Set*

$$T_0 = \left\lceil \frac{2}{\alpha\eta_0} \log\left(\frac{2R}{\varepsilon}\right) \right\rceil, \quad T_k = \left\lceil \frac{2\log(4)}{\alpha\eta_k} \right\rceil \quad \text{for } k \geq 1, \quad \text{and} \quad K = \left\lceil 1 + \log_2\left(\frac{2\eta_0\sigma^2}{\alpha\varepsilon}\right) \right\rceil.$$

The final iterate x produced satisfies $\mathbb{E} \|x - x^\|^2 \leq \varepsilon$, while the total number of iterations of stochastic gradient play called is at most*

$$\mathcal{O}\left(\frac{L^2}{\alpha^2} \cdot \log\left(\frac{2R}{\varepsilon}\right) + \frac{\sigma^2}{\alpha^2\varepsilon}\right).$$

6.3 Adaptive Gradient Method in Performative Prediction Games

Throughout this section, we continue working under the parametric Assumption 6. An apparent deficiency of the stochastic gradient method discussed in Section 6.2 is that each player i needs to know the matrix A_i that governs the performative effect of the player on the distribution. In typical settings, the matrix A_i may be unknown to the player, but it might be possible to estimate it from data. Inspired by methods in adaptive control to simultaneously estimate the parameters of the system and optimize the control input, we propose the *adaptive gradient method* outlined in Algorithm 1.⁵ In each iteration, each player simultaneously estimates their distribution parameters and myopically optimizes their individual loss via stochastic gradient method on the current estimated loss. More precisely, the algorithm maintains two sequences: (i) x^t that eventually converges to the Nash equilibrium x^* , and (ii) estimates $\hat{A}_i^t$ that dynamically estimates the unknown matrix $\bar{A}_i$. In each iteration t , the algorithm draws samples $z_i^t \sim \mathcal{D}_i(x^t)$, and each player i takes the gradient step

$$x_i^{t+1} = \text{proj}_{\mathcal{X}_i} \left(x_i^t - \eta_t ((\nabla_i \ell_i(x^t, z_i^t) + (\hat{A}_{ii}^t)^\top \nabla_{z_i} \ell_i(x^t, z_i^t))) \right),$$

where $\hat{A}_{ii}^t$ denotes the submatrix of $\hat{A}_i^t$ whose columns are indexed by player i 's action space. Next, in order to update $\hat{A}^t$, the algorithm draws a sample $q_i^t \sim \mathcal{D}_i(x^t + u^t)$ where u^t is a user-specified noise sequence. Observe that conditioned on u^t , the equality holds:

$$\mathbb{E}[q_i^t - z_i^t \mid u^t] = \bar{A}_i u^t.$$

Therefore, a good strategy for forming a new estimate $\hat{A}_i^{t+1}$ of $\bar{A}_i$ from $\hat{A}_i^t$ is to take a gradient step on the least squares objective

$$\min_{B_i} \frac{1}{2} \|q_i^t - z_i^t - B_i u^t\|^2.$$

5. We remark that the word “adaptive” here refers to adaptively estimating the model parameters, and is different from its meaning in methods like AdaGrad, where it is the algorithm’s stepsize that is being adapted.

In particular, a key step in the adaptive gradient method is running online least squares in parallel with the stochastic gradient play (cf. step 7 in Algorithm 1). Explicitly, this gives the update

$$\hat{A}_i^{t+1} = \hat{A}_i^t + \nu_t(q_i^t - z_i^t - \hat{A}_i^t u^t)(u^t)^\top.$$

Analogous to estimation in adaptive control (Varaiya and Wets, 1988) or machine learning, we exploit noise injection u_t to ensure sufficient exploration of the parameter space. In particular, the noise vector needs to be sufficiently isotropic. This ensures that the covariance of the estimates of the parameters is bounded away from zero. The idea is that it is possible that the iterates of the algorithm would not necessarily vary in enough directions for the least squares solution to converge – by injecting noise, we can guarantee that we can estimate each entry of the matrix A_i even in the worst case. We impose the following assumption.

Assumption 9 (Injected Noise) The injected noise vector $u^t = (u_1^t, \dots, u_n^t) \in \mathbb{R}^d$ is a zero-mean random vector that is independent of x^t , and independent of the injected noise at any previous queries to the environment by any player. Moreover, there exists constants $c_l, R > 0$ and $c_{u,i} > 0$ for each $i \in [n]$ such that for all $t \geq 0$ and $i \in [n]$ the random vector $v_i := u_i^t$ satisfies

$$0 \prec c_l \cdot I \preceq \mathbb{E}[v_i v_i^\top], \quad \mathbb{E} \|v_i\|^2 \leq c_{u,i}, \quad \text{and} \quad \mathbb{E}[\|v_i\|^2 v_i v_i^\top] \preceq R^2 \mathbb{E}[v_i v_i^\top].$$

In the simple Gaussian case where $u_t \sim \mathcal{N}(0, I_d)$, we may set⁶

$$c_l = 1, \quad c_{u,i} = d_i, \quad \text{and} \quad R^2 = 3 \max_{i \in [n]} d_i.$$

Analyzing the convergence of Algorithm 1 amounts to decomposing the analysis into convergence of the stochastic gradient method on the estimated losses induced by the sequence of $\hat{A}_i^t$, and convergence of the estimation error $\mathbb{E} \|\hat{A}_i^t - \bar{A}_i\|^2$. The former analysis proceeds in an analogous fashion to that of Theorem 19 in Section 6.2. For the latter, we leverage the injected noise to ensure there is sufficient *exploration*. The following lemma establishes a one-step improvement guarantee on estimation of $\bar{A}_i$. Throughout, we set $\hat{A}^t := (\hat{A}_1^t, \dots, \hat{A}_n^t)$ and let $\|\cdot\|_F$ denote the Frobenius norm. We also let $\mathbb{E}_t$ be the conditional expectation with respect to the σ -algebra generated by $(x^l, u^l)_{l=1, \dots, t}$.

Lemma 21 (Estimation error) *Suppose that Assumptions 6 and 9 hold and choose $\nu_t \in (0, \frac{2}{R^2})$. Then the matrices $\hat{A}_i^t$ generated by Algorithm 1 satisfy the estimate:*

$$\frac{1}{2} \mathbb{E}_t \|\hat{A}_i^{t+1} - \bar{A}_i\|_F^2 \leq \frac{1 - c_l \nu_t (2 - \nu_t R^2)}{2} \|\hat{A}_i^t - \bar{A}_i\|_F^2 + \nu_t^2 \text{tr}(\Sigma_i) c_{u,i}. \quad (34)$$

Therefore when setting $\nu_t = \frac{2}{(c_l(t + \frac{2R^2}{c_l}))}$ for all $t \geq 0$, the estimate holds:

$$\mathbb{E} \|\hat{A}^t - \bar{A}\|_F^2 \leq \frac{\max \left\{ (1 + \frac{2R^2}{c_l}) \|\hat{A}_1 - \bar{A}\|_F^2, \frac{8 \sum_{i=1}^n \text{tr}(\Sigma_i) c_{u,i}}{c_l^2} \right\}}{t + \frac{2R^2}{c_l}}.$$

6. For the justification of the expression for R^2 , see (Dieuleveut et al., 2017, Section 2.1).

Algorithm 1: Adaptive Gradient Method

```

1 Input: Stepsizes  $\{\eta_t\}_{t \geq 1}$ ,  $\{\nu_t\}_{t \geq 1}$ ; initial  $x^1 \in \mathbb{R}^d$ ,  $\hat{A}_i^1 \in \mathbb{R}^{m \times d}$ ;
2 for  $t = 1, \dots, T$  do
3     for  $i \in [n]$  do
4         Query the environment: Draw samples  $z_i^t \sim \mathcal{D}_i(x^t)$  and  $q_i^t \sim \mathcal{D}_i(x^t + u^t)$ ;
5         Individual gradient update:
            
$$x_i^{t+1} = \text{proj}_{\mathcal{X}_i} \left( x_i^t - \eta_t (\nabla_i \ell_i(x^t, z_i^t) + (\hat{A}_{ii}^t)^\top \nabla_{z_i} \ell_i(x^t, z_i^t)) \right),$$

6         where  $\hat{A}_{ii}^t$  denotes the submatrix of  $\hat{A}_i^t$  whose columns are indexed by player  $i$ .
7         Estimation update:  $\hat{A}_i^{t+1} = \hat{A}_i^t + \nu_t (q_i^t - z_i^t - \hat{A}_i^t u_i^t)(u_i^t)^\top$ 
8     end
9 end
    
```

Proof This follows from a standard estimate for online least squares, which appears as Lemma 27 in Appendix C. Namely, let $\mathcal{G}_1$ be the σ -algebra generated by $x^1, \dots, x^t$ and let $\mathcal{G}_2$ be the σ -algebra generated by $\mathcal{G}_1 \cup \{u^t\}$. Set $b = q_i^t - z_i^t$, $y = u_i^t$, $B = \hat{A}_i^t$, $V = \bar{A}_i$, $v = u_i^t$, $\lambda_1 = c_l$, and $\lambda_2 = c_{u,i}$.

Let us upper bound the variance $\mathbb{E}[\|Vy - b\|^2 \mid \mathcal{G}_2]$. To this end, let w and w' be drawn i.i.d from $\mathcal{P}_i$. Observe that conditioned on u_i^t , the random vector $\bar{A}_i u_i^t - (q_i^t - z_i^t)$ has the same distribution as $w - w'$. Let us compute

$$\mathbb{E} \|w - w'\|^2 = \text{tr}(\mathbb{E}((w - w')(w - w')^\top)) = 2\text{tr}(\Sigma_i).$$

Therefore, we may set $\sigma^2 = 2\text{tr}(\Sigma_i)$. An application of Lemma 27 completes the proof of (34). Summing up (34) for $i = 1, \dots, n$ and using the tower rule for conditional expectations yields:

$$\mathbb{E} \|\hat{A}^{t+1} - \bar{A}\|_F^2 \leq (1 - \nu_t c_l (2 - \nu_t^2 R^2)) \mathbb{E} \|\hat{A}^t - \bar{A}\|_F^2 + 2\nu_t^2 \sum_{i=1}^n \text{tr}(\Sigma_i) c_{u,i}.$$

Noting $\nu_t \leq \frac{1}{R^2}$, we deduce $1 - \nu_t c_l (2 - \nu_t^2 R^2) \leq 1 - \nu_t c_l$. The result follows directly from plugging in the value of ν_t and using Lemma 25 in Appendix B. $\blacksquare$

Next we show that the direction of motion of Algorithm 6.3 is well-aligned with the direction of motion of the stochastic gradient method. To this end, define the true (stochastic) vector of individual gradients

$$v^t := (\nabla_i \ell_i(x^t, z_i^t) + A_i^\top \nabla_{z_i} \ell_i(x^t, z_i^t))_{i \in [n]},$$

and its estimator that is used by the algorithm

$$\hat{v}^t := (\nabla_i \ell_i(x^t, z_i^t) + (\hat{A}_{ii}^t)^\top \nabla_{z_i} \ell_i(x^t, z_i^t))_{i \in [n]}.$$

We make the following Lipschitzness assumption on the loss $\ell_i(x, z_i)$ in the variable z_i .

Assumption 10 (Lipschitz continuity in z) Suppose that there exists a constant $\delta > 0$ such that for all $x \in \mathcal{X}$, the estimate holds:

$$\mathbb{E}_{z \sim \mathcal{D}_\pi(x)} \sqrt{\sum_{i=1}^n \|\nabla \ell_i(x, z_i)\|^2} \leq \delta.$$

Lemma 22 Suppose that Assumptions 6 and 10 hold. Then for each $t \geq 1$ and $i \in [n]$, the estimate holds:

$$\mathbb{E}_t \|\hat{v}^t - v^t\| \leq \delta \|\hat{A}^t - \bar{A}\|_F^2.$$

Proof Notice that we may write $\hat{v}^t - v^t = B^t w^t$, where B^t is the block diagonal matrix with blocks $\hat{A}_{ii}^t - A_i$ and we set $w^t = (\nabla_{z_i} \ell_i(x^t, z_i^t))_{i=1}^n$. Using Hölder's inequality, we obtain the following estimate as claimed:

$$\mathbb{E}_t \|\hat{v}^t - v^t\| = \mathbb{E}_t \|B^t w^t\| \leq \|B^t\|_F \cdot \mathbb{E}_t \|w^t\| \leq \delta \|\hat{A}^t - \bar{A}\|_F^2.$$

■

In light of Lemmas 21 and 22, we may interpret Algorithm 6.3 as an approximation to the stochastic gradient method with a bias that tends to zero; we may then simply invoke generic convergence guarantees for biased stochastic gradient methods, which we record in Theorem 24 of Appendix A. We will make use of the following assumption.

Assumption 11 (Finite variance) Suppose that there exists $\sigma > 0$ such that for all $x \in \mathcal{X}$, the variance bound holds:

$$\mathbb{E}_{z_i \sim \mathcal{D}_i(x^t)} \|\nabla_{i, z_i} \ell_i(x^t, z_i^t) - \mathbb{E}_{z'_i \sim \mathcal{D}_i(x^t)} \nabla_{i, z_i} \ell_i(x^t, z'_i)\|^2 \leq \sigma^2.$$

The end result is the following theorem, which in particular implies a $\mathcal{O}(d/t)$ rate of convergence when u^t are standard Gaussian. See the discussion after the theorem.

Theorem 23 (Convergence of the adaptive method) Suppose that Assumptions 6, 7, 9, 10, and 11 hold and that the game (4) is α -strongly monotone. Define the constant $k_0 = 1 + \frac{8L^2}{\alpha^2}$ and $q_0 = \frac{2R^2}{c_l}$ and set $\eta_t = \frac{2}{\alpha(t+k_0-2)}$ and $\nu_t = \frac{2}{c_l(t+q_0)}$ for all $t \geq 0$. Then for all $t \geq 1$, the iterates generated by Algorithm 1 satisfy

$$\begin{aligned} \mathbb{E} \|x^t - x^*\|^2 \leq & \frac{\max \left\{ \frac{1}{2} \alpha^2 (1 + k_0) \|x_1 - x^*\|^2, \ 32(1 + 2\|\bar{A}\|_F^2) \sigma^2 + 8\delta^2 Z \max \left\{ \frac{1+k_0}{1+q_0}, 1 \right\} \right\}}{\alpha^2(t + k_0)} \\ & + \frac{\max \left\{ \frac{1}{2} \alpha^2 (1 + k_0)^{3/2} \|x_1 - x^*\|^2, \ 64\sigma^2 Z \max \left\{ \frac{1+k_0}{1+q_0}, 1 \right\} \right\}}{\alpha^2(t + k_0)^{3/2}}. \end{aligned}$$

where we set $Z = \max \left\{ (1 + \frac{2R^2}{c_l}) \|\hat{A}^1 - \bar{A}\|_F^2, \frac{8 \sum_{i=1}^n \text{tr}(\Sigma_i) c_{u,i}}{c_l^2} \right\}$.

Proof We will apply the standard convergence guarantees in Theorem 24 of Appendix A for biased stochastic gradient methods. Using Lemma 22 we estimate the gradient bias:

$$\|\mathbb{E}_t[\hat{v}^t] - \mathbb{E}_t[v^t]\| = \mathbb{E}_t\|\hat{v}^t - v^t\| \leq \delta\|\hat{A}^t - \bar{A}\|_F^2.$$

Next, we estimate the variance:

$$\mathbb{E}_t[\|\hat{v}_i^t - \mathbb{E}\hat{v}_i^t\|^2] = \mathbb{E}_t\left\|\begin{bmatrix} I & 0 \\ 0 & \hat{A}_{ii}^t \end{bmatrix} (\nabla_{i,z_i} \ell_i(x^t, z_i^t) - \mathbb{E}_{z'_i \sim \mathcal{D}_i(x^t)} \nabla_{i,z_i} \ell_i(x^t, z'_i))\right\|^2.$$

Summing these inequalities over $i \in [n]$, we deduce

$$\mathbb{E}[\|\hat{v}_i^t - \mathbb{E}\hat{v}_i^t\|^2] \leq \max\{1, \|\hat{A}^t\|_{\text{op}}^2\} \sigma^2.$$

Recalling the definition of Z and q_0 and applying Theorem 24 in Appendix A we deduce

$$\begin{aligned} \mathbb{E}_t\|x^{t+1} - x^\star\|^2 &\leq \frac{1}{1 + \eta_t \alpha} \|x^t - x^\star\|^2 + \frac{2\eta_t^2 (\max\{1, \|\hat{A}^t\|_{\text{op}}^2\}) \sigma^2}{1 + \eta_t \alpha} + \frac{2\eta_t \delta^2 \|\hat{A}^t - \bar{A}\|_F^2}{\alpha} \\ &\leq \frac{1}{1 + \eta_t \alpha} \|x^t - x^\star\|^2 + 2\eta_t^2 (1 + \|\hat{A}^t\|_F^2) \sigma^2 + \frac{2\eta_t \delta^2 \|\hat{A}^t - \bar{A}\|_F^2}{\alpha} \\ &\leq \frac{1}{1 + \eta_t \alpha} \|x^t - x^\star\|^2 + 2\eta_t^2 (1 + 2\|\bar{A}\|_F^2) \sigma^2 + \frac{2\eta_t \delta^2 \|\hat{A}^t - \bar{A}\|_F^2}{\alpha} \\ &\quad + 4\eta_t^2 \sigma^2 \|\hat{A}^t - \bar{A}\|_F^2, \end{aligned}$$

where the last inequality follows from the algebraic expression $\|\hat{A}^t\|^2 \leq 2\|\bar{A}\|^2 + 2\|\hat{A}^t - \bar{A}\|^2$. Taking expectations and using the tower rule, we compute

$$\begin{aligned} \mathbb{E}\|x^{t+1} - x^\star\|^2 &\leq \frac{1}{1 + \eta_t \alpha} \mathbb{E}\|x^t - x^\star\|^2 + 2\eta_t^2 (1 + 2\|\bar{A}\|_F^2) \sigma^2 + \frac{2\eta_t \delta^2 \mathbb{E}\|\hat{A}^t - \bar{A}\|_F^2}{\alpha} \\ &\quad + 4\eta_t^2 \sigma^2 \mathbb{E}\|\hat{A}^t - \bar{A}\|_F^2 \\ &\leq \frac{1}{1 + \eta_t \alpha} \mathbb{E}\|x^t - x^\star\|^2 + 2\eta_t^2 (1 + 2\|\bar{A}\|_F^2) \sigma^2 + \frac{2\eta_t \delta^2 Z}{\alpha(t + q_0)} + \frac{4\eta_t^2 \sigma^2 Z}{t + q_0}, \end{aligned}$$

where the last inequality follows from Lemma 21.

Now our choice $\eta_t = \frac{2}{\alpha(t+k_0-2)}$ ensures $\frac{1}{1+\eta_t\alpha} = 1 - \frac{2}{t+k_0}$. Therefore we deduce

$$\begin{aligned} \mathbb{E}\|x^{t+1} - x^\star\|^2 &\leq \left(1 - \frac{2}{t+k_0}\right) \mathbb{E}\|x^t - x^\star\|^2 + \frac{8(1 + 2\|\bar{A}\|_F^2) \sigma^2}{\alpha^2(t+k_0-2)^2} \\ &\quad + \frac{16\sigma^2 Z}{\alpha^2(t+q_0)(t+k_0-2)^2} + \frac{4\delta^2 Z}{\alpha^2(t+q_0)(t+k_0-2)}. \end{aligned} \tag{35}$$

We now aim to apply Lemma 26 in Section B. To this end, we need to upper bound the last three terms in (35) so that the denominators are of the form $(t+k_0)^p$ for some power p . To this end, note the following elementary estimates:

$$\frac{t+k_0}{t+k_0-2} \leq \frac{k_0+1}{k_0-1} \leq 2$$

and

$$\frac{(t+k_0)^2}{(t+q_0)(t+k_0-2)} \leq \frac{k_0+1}{k_0-1} \cdot \frac{t+k_0}{t+q_0} \leq \frac{c(k_0+1)}{k_0-1} \leq 2c$$

where $c = \max_{t \geq 1} \{\frac{t+k_0}{t+q_0}\} = \max\{\frac{1+k_0}{1+q_0}, 1\}$. Combining these estimates with (35), we obtain

$$\mathbb{E}\|x^{t+1} - x^*\|^2 \leq \left(1 - \frac{2}{t+k_0}\right) \|x^t - x^*\|^2 + \frac{32(1+2\|\bar{A}\|_F^2)\sigma^2 + 8\delta^2 Zc}{\alpha^2(t+k_0)^2} + \frac{64\sigma^2 Z \cdot c}{(\alpha^2(t+k_0)^3)}.$$

Applying Lemma 26 in Section B, we conclude:

$$\begin{aligned} \mathbb{E}\|x^t - x^*\|^2 &\leq \frac{\max\left\{\frac{1}{2}\alpha^2(1+k_0)\|x_1 - x^*\|^2, 32(1+2\|\bar{A}\|_F^2)\sigma^2 + 8\delta^2 Zc\right\}}{\alpha^2(t+k_0)} \\ &\quad + \frac{\max\left\{\frac{1}{2}\alpha^2(1+k_0)^{3/2}\|x_1 - x^*\|^2, 64\sigma^2 Zc\right\}}{\alpha^2(t+k_0)^{3/2}}. \end{aligned}$$

This completes the proof. ■

In particular, consider the Gaussian case $u^t \sim \mathcal{N}(0, I_d)$ in the setting when $d_i = C_i d$ for some numerical constants C_i , and when the traces $\text{tr}(\Sigma_i) \equiv \text{tr}(\Sigma)$ are equal for all $i \in [n]$. Then the efficiency estimate in Theorem 23 becomes

$$\begin{aligned} \mathbb{E}\|x^t - x^*\|^2 &= \mathcal{O}\left(\frac{\max\left\{L^2\|x_1 - x^*\|^2, \|\bar{A}\|_F^2\sigma^2 + \delta^2 \max\{d\|\hat{A}^1 - \bar{A}\|_F^2, \text{tr}(\Sigma) \sum_{i=1}^n c_{u,i}\} \max\{\frac{L^2}{d\alpha^2}, 1\}\right\}}{\alpha^2 t + L^2}\right. \\ &\quad \left. + \frac{\max\left\{L^3\|x_1 - x^*\|^2, \alpha\sigma^2 \max\{d\|\hat{A}^1 - \bar{A}\|_F^2, \text{tr}(\Sigma) \sum_{i=1}^n c_{u,i}\} \max\{\frac{L^2}{d\alpha^2}, 1\}\right\}}{(\alpha^2 t + L^2)^{3/2}}\right). \end{aligned}$$

Consequently, treating all terms besides d and t as constants, yields the rate $\mathcal{O}(\frac{d}{t})$.

7. Numerical Examples

Section 5 introduced two examples that are motivated by practical problems: revenue maximization in ride-share markets and interactions between election prediction platforms. In this section, to illustrate the theoretical results we explore a synthetic multiplayer performative prediction game which is an abstraction of the latter, and a semi-synthetic ride-share market example. For brevity, the majority of the numerical experiments on the revenue maximization in ride-share markets using real data are relegated to Appendix E.

7.1 Multiplayer Performative Prediction with Strategic Sources

The purpose of the synthetic example presented in this section is to illustrate the theoretical convergence results for the algorithms proposed in Section 6, including the setting in which players receive feedback asynchronously.⁷

7. For the repeated stochastic gradient method analyzed in Section 4, we include a plot confirming the theoretical convergence guarantee in Appendix E.

Game Abstraction. In Example 2, we introduced a performative prediction game motivated by multiple election platforms. A set of features θ are drawn i.i.d. from a static distribution $\mathcal{P}_\theta$, and player i observes a sample drawn from the conditional distribution

$$z_i|\theta \sim \varphi_i(\theta) + A_i x_i + A_{-i} x_{-i} + w_i,$$

where $\varphi_i : \mathbb{R}^{d_i} \rightarrow \mathbb{R}^{m_i}$ is an arbitrary map, the parameter matrices $A_i \in \mathbb{R}^{m_i \times d_i}$ and $A_{-i} \in \mathbb{R}^{m_i \times d_{-i}}$ are fixed, and w_i is a zero-mean random variable with variance $\sigma_{w_i}^2$. Each player seeks to predict z_i by minimizing the loss

$$\ell_i(x_i, z_i) = \frac{1}{2} \|z_i - \theta^\top x_i\|^2.$$

We showed in Example 2 that the map $x \mapsto H_x(y)$ is monotone as long as

$$\sqrt{n-1} \cdot \max_{i \in [n]} \|A_{-i}^\top A_i\|_{\text{op}} \leq \min_{i \in [n]} \lambda_{\min}(A_i^\top A_i), .$$

As noted in Example 2, the interpretation is that the performative effects due to interaction with competitors are small relative to any player’s own performative effects. We enforce this condition on the parameters for our numerical examples.

Instance generation. We randomly generate problem instances—namely, the parameters A_i, A_{-i} for $i \in [n]$ —by using `scipy.sparse.random` which allows for the sparsity of the matrix to be set in addition to randomly generating the matrix parameters.⁸ Furthermore, we set $\theta \in \mathbb{R}^{d \times m}$ with entries distributed as $\mathcal{N}(0, 0.01)$, $\sigma_w^2 = 0.1$, and $\varphi_i(\theta) = \theta^\top \mathbf{1}_{d \times 1}$ for $d = 2$ and $m = 5$ for the experiments in Figure 1. These values can be changed in the provided code, resulting in similar conclusions regarding the convergence rate.

As discussed in earlier sections, in many practical applications players may not receive information synchronously. To model this, we consider the setting in which player i receives information in each iteration with probability p_i .

Experiment 1a: Convergence rates & Estimation Error. In Figure 1a we show the iteration complexity of the norm-square error to the Nash equilibrium for the stochastic gradient method, the adaptive gradient method, and players independently playing according to derivative free optimization.⁹ Figure 1b shows the estimation error for the matrices (A_i, A_{-i}) , $i = 1, 2$ in the adaptive gradient method. We run the methods for 20 different random initializations and show the mean and ± 1 standard deviation, depicted with darker lines, and the individual sample trajectories using lighter shade lines of the same color indicated in the legend. The step-sizes are chosen according to the theory.

Experiment 1b: Asynchronous updates. Figure 2a contains a violin plot for the number of iterations required to hit $\varepsilon = 1\text{e-}3$ error (to the Nash equilibrium) in the asynchronous setting where $p_1 = 1$ and we vary p_2 , as shown on the x -axis. As expected, we see that as p_2 increases from zero to one, the number of iterations required decreases. Additionally, as expected, the distributions for adaptive gradient play have a large spread due to the time it takes players to estimate the performative effects.

8. The data and code used in this paper are publicly available (<https://github.com/ratliff1j/performativepredictiongames>).

9. In all the experiments, we compute the Nash equilibrium for the game defined by the expected cost using a symbolic solver such as Mathematica or `sympy`.

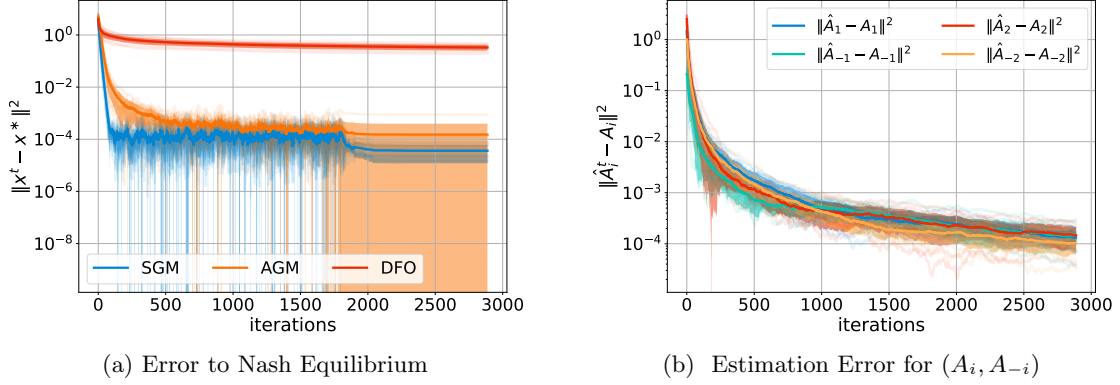


Figure 1: **Strategic Prediction.** (a) Iteration complexity of the norm-squared error to the Nash equilibrium for both the stochastic gradient method (SGM), adaptive gradient method (AGM), and derivative free optimization (DFO) for a randomly generated problem instance where $m_i = 5$ and $d_i = 2$ for each $i \in \{1, 2\}$. The sample complexity of AGM empirically matches that of SGM as expected up to the estimation error. (b) Estimation error for the matrices (A_i, A_{-i}) , $i = 1, 2$. Due to the time it takes to estimate the matrices we see in (a) that AGM converges somewhat slower.

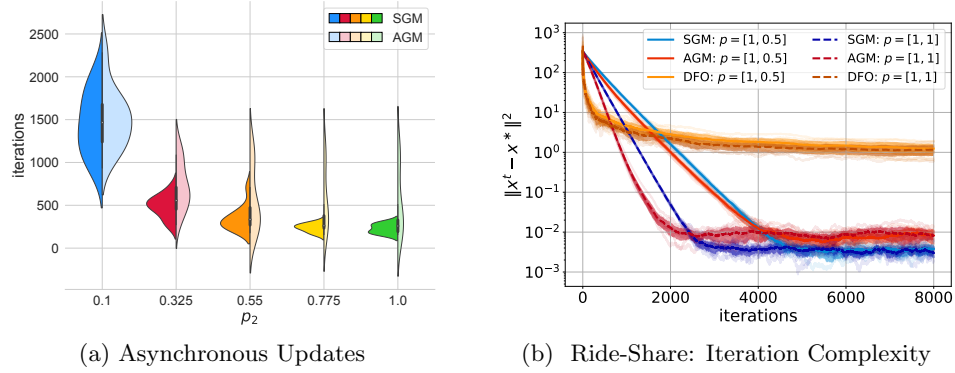


Figure 2: (a) **Asynchronous Updates in Multiplayer Strategic Prediction (§7.1):** Number of iterations to hit an ε -Nash equilibrium where $\varepsilon = 1e-3$ for stochastic gradient play (SGM) and adaptive gradient play (AGM), respectively, where $p = (p_1, p_2)$ with $p_1 = 1$ and p_2 varying from 0.1 to one. (b) **Competition in Ride-Share Markets (§7.2):** Iteration complexity for Nash seeking algorithms including derivative free optimization(DFO), stochastic gradient method and adaptive gradient method.

7.2 Revenue Maximization in Ride-Share Markets

In this section, we explore semi-synthetic competition between two ride-share platforms seeking to maximize their revenue given that the demand they experience is influenced by

their own prices as well as their competitors. We use data from a prior Kaggle competition to set up the semi-synthetic simulation environment.¹⁰

Game Abstraction. The abstraction for the game can be described as follows. Consider a ride-share market with two platforms that each seek to maximize their revenue by setting prices for their rides given by a vector $x_i \in \mathbb{R}^{m_i}$. The vector of demands $z_i \in \mathbb{R}^{m_i}$ containing demand information for m_i locations that each ride-share platform serves is influenced not only by the prices they set but also the prices that their competitor sets. Suppose that platform i 's loss is given by

$$\ell_i(x_i, z_i) = -\frac{1}{2}z_i^\top x_i + \frac{\lambda_i}{2}\|x_i\|^2$$

where $\lambda_i \geq 0$ is some regularization parameter, and $x_i \in \mathbb{R}^{m_i}$ represents the vector of price differentials from some nominal price for each of the m_i locations. Observe that this game is separable since the losses ℓ_i do not explicitly depend on x_{-i} . Moreover, let us suppose that the random demand z_i takes the semi-parametric form $z_i = \zeta_i + A_i x_i + A_{-i} x_{-i}$, where ζ_i follows some base distribution $\mathcal{P}_i$ and the parameters A_i and A_{-i} capture price elasticities to player i 's and its competitor's change in price, respectively. We have that $A_i \preceq 0$ and $A_{-i} \succeq 0$; this captures that an increase in player i 's prices results in decreased demand, while an increase in its competitor's prices results in increased demand. Moreover, we showed in Example 1 that the mapping $x \mapsto H_x(y)$ is trivially monotone. Hence, the game between ride-share platforms is strongly monotone and admits a unique Nash equilibrium. Throughout the remainder of this section, we set $\lambda_1 = \lambda_2 = 1$.

Instance Generation. There are eleven locations, and each element in x_i represents the price difference (set by platform i) from a nominal price at each location. We aggregate the rides into bins of \$5 increments; this is done by taking the raw data and rounding the price to the nearest bin as follows: $p_{\text{nominal}} = 5 \cdot \lfloor \frac{p}{5} \rfloor$ where p is the actual price of a particular ride. Then, for each bin j we have an empirical distribution $\mathcal{P}_{i,j}$ for each player $i \in \{1, 2\}$ which is just the collection of rides in the data set for the location and price range specified by that bin. The results presented in Figure 2b use data for the \$10 nominal price bin; however, in the linked code it is easy to adjust this parameter to any of the other price bins, and the conclusions we draw are similar across the bins. The results in Figure 3 use data from the bins ranging from \$10 to \$30 increments of \$5.

In the experiments presented, we estimate the matrices A_i and A_{-i} that govern the performative effects from the data. The details of this estimation are outlined in detail Appendix E, and the heuristics used to set-up the semi-synthetic model can be changed in the code-base.¹¹ The semi-synthetic model is constructed such that there are no performative effects across different locations. This amounts to zero elements off the diagonals of the matrices A_i and A_{-i} . However, in the provided code base, we have additional experiments that estimate the correlation between locations and explore the effects of positive and negative correlations on equilibrium outcomes when the off-diagonal elements are nonzero.

We run each of the algorithms in Section 6 from twenty random initial conditions, and compute the error between the trajectory of the algorithm and the Nash equilibrium. The parameters used in the algorithms are set based on the respective theorems.

10. The data used in this paper is publicly available (<https://www.kaggle.com/br11rb/uber-and-lyft-dataset-boston-ma>).

11. In our experience, changing the heuristics does not significantly change the trends we observe.

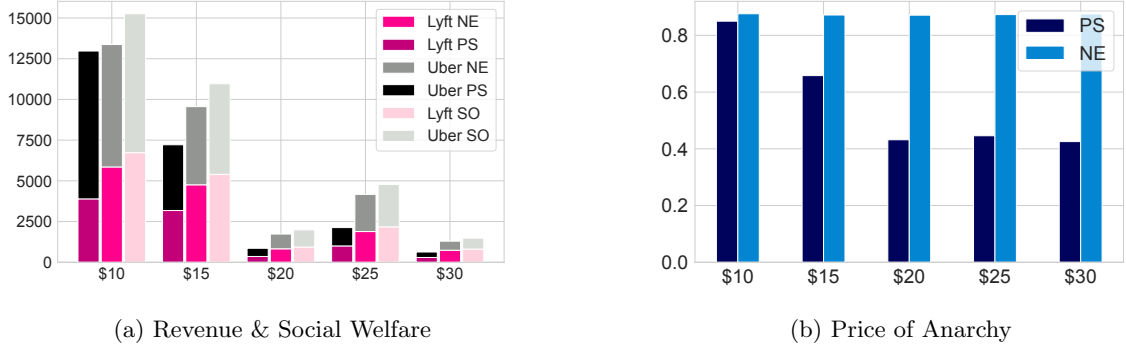


Figure 3: **Competition in Ride-Share Markets:** (a) Revenue and social welfare for each of the bins ranging from \$10 to \$30 increments of \$5. The “grey” shaded parts of the bar show Uber’s revenue while the “pink” shaded parts show Lyft’s. The sum of the two is the social welfare. (b) Price of Anarchy for each of the bins ranging from \$10 to \$30 increments of \$5.

Below we describe two experiments: (2a) verifies convergence rates, and (2b) explores the relative social cost and market share split under different equilibrium concepts. Further numerical experiments exploring the comparison between the social cost at the different equilibrium concepts, as well as comparing outcomes when one player ignores performative effects while the other does not are contained in Appendix E, and theoretical bounds on the suboptimality gaps in Appendix D.

Experiment 2a: Convergence Rates. Figure 2b simply illustrates the convergence rate as predicted by the theory in Section 6 for each of the algorithms therein. We show the mean of the error trajectories and ± 1 one standard deviation. The plots demonstrate that the empirical sample complexity of the adaptive gradient method and the stochastic gradient method are nearly identical, and outperform the derivative free method as expected.

Experiment 2b: Social Efficiency of Different Equilibrium Concepts. As noted in the preceding sections, in the study of equilibrium for games, it is important to understand the efficiency of different equilibrium concepts. The typical benchmark for efficiency is the cost at the social optimum. The social cost $\mathcal{C}$ is defined as the sum of all the players individual costs, and the negative of this term is the social welfare $\mathcal{S}$:

$$\mathcal{C}(x) = \sum_{i=1}^n \mathcal{L}_i(x) \quad \text{and} \quad \mathcal{S} \equiv -\mathcal{C}.$$

Let x^{ne} be the Nash equilibrium of the game $(\mathcal{L}_1, \mathcal{L}_2)$ and let x^{ps} be the Nash equilibrium of the game $\mathcal{G}(x^{\text{ps}})$, using the notation from Section 3 for the game induced by x^{ps} . We find the unique socially optimal equilibrium x^{so} (the social cost is strongly convex), and the Nash and performatively stable equilibrium using the expected cost in a symbolic solver (Mathematica).

The price anarchy is a common metric for equilibrium efficiency and is defined as the ratio of the social welfare at the worst case competitive equilibrium concept relative to the social welfare at the social optimum—namely, it is given by

$$\text{PoA}(x) = \frac{\max_{x \in \text{Eq}(\bar{\mathcal{G}})} \mathcal{S}(x)}{\mathcal{S}(x^{\text{so}})},$$

where $\text{Eq}(\bar{\mathcal{G}})$ denotes the set of equilibria for the game $\bar{\mathcal{G}} = (\mathcal{L}_1, \dots, \mathcal{L}_n)$. An equilibrium concept is said to be more efficient the closer this quantity is to one.

In Figure 3, we show the (a) revenue and the (b) price of anarchy in each of the price bins ranging from \$10 to \$30 in increments of \$5. It is interesting to note the price of anarchy of the Nash equilibrium for each of the price bins is roughly the same (~ 0.87) while the price of anarchy for the performatively stable equilibrium decreases as the base bin value increases. Looking at the revenue, the social cost (i.e. the total height of the bar) is higher in all of the bins. Lyft makes progressively more moving from performatively stable (PS) to Nash (NE) to the social optimum (SO) thereby indicating that taking into account performative effects is good for Lyft. This is in part because Lyft is the “smaller” player: since there is less demand for Lyft on average than Uber in the base demand $\mathcal{P}_i$, accounting for performative effects is strategically advantageous. Uber also does marginally better moving from PS to NE to SO, however the gain is smaller than that of its competitor. This suggests investigating how market power (size of the base market) plays a role in whether or not firms should invest in accounting for performativity, versus the simpler repeated retraining process.

In the same vein as Lemma 6, in Appendix D we have included theoretical results bounding the gaps between the Nash equilibrium, the performatively stable equilibrium, the social optimum, and the stable equilibrium reached when some players run the repeated stochastic gradient method (Section 4) and others run a Nash seeking algorithm (Section 6). We use these bounds to obtain bounds on the suboptimality gaps for players relative to the social optimum. Numerical results exploring these gaps, and the impact on the losses experience by players are incorporated in Appendix E.

8. Discussion

The new class of games in this paper motivates interesting future work at the intersection of statistical learning theory, optimization, and game theory. For instance, it is of interest to extend the present framework to handle more general parametric forms of the distributions $\mathcal{D}_i$. Many multiplayer performative prediction problems exhibit a hierarchical structure such as a governing body that presides over an institution; hence, a Stackelberg variant of multiplayer performative prediction is of interest. Along these lines, the multiplayer performative prediction problem is also of interest for mechanism design problems arising in applications such as recommender systems. For instance, the recommendations that platforms select at the Nash equilibrium influence the preferences of consumers (data-generators). A mechanism designer (e.g., the government) can place constraints on platforms to prevent them from *manipulating* users’ preferences to make their prediction tasks easier.

Appendix A. Stochastic gradient method with bias

In this section, we consider a variational inequality

$$0 \in G(x) + N_{\mathcal{X}}(x), \quad (36)$$

where $\mathcal{X} \subset \mathbb{R}^d$ is a closed convex set and $G: \mathbb{R}^d \rightarrow \mathbb{R}^d$ is an L -Lipschitz continuous and α -strongly monotone map. We will analyze the stochastic gradient method, which in each iteration performs the update:

$$x^{t+1} = \text{proj}_{\mathcal{X}}(x^t - \eta g^t), \quad (37)$$

where $\eta > 0$ is a fixed stepsize and g_t is a sequence of random variables, which approximates $G(x^t)$. In particular, it will be crucial for us to allow g^t to be a *biased* estimator of $G(x^t)$. Formally, we make the following assumption on the randomness in the process. Throughout, x^* denotes the unique solution of (36).

Assumption 12 (Stochastic framework) *Suppose that there exists a filtered probability space $(\Omega, \mathcal{F}, \mathbb{F}, \mathbb{P})$ with filtration $\mathbb{F} = (\mathcal{F}_t)_{t \geq 0}$ such that $\mathcal{F}_0 = \{\emptyset, \Omega\}$. Suppose moreover that g_t is $\mathcal{F}_{t+1}$ -measurable and there exist constants $B, \sigma \geq 0$ and $\mathcal{F}_t$ -measurable random variables $C_t, \sigma_t \geq 0$ satisfying the bias/variance bounds*

$$\begin{aligned} \text{(Bias)} \quad & \|\mathbb{E}_t g^t - G(x^t)\| \leq C_t + B\|x^t - x^*\|, \\ \text{(Variance)} \quad & \mathbb{E}_t \|g^t - \mathbb{E}_t[g^t]\|^2 \leq \sigma_t^2 + D^2\|x^t - x^*\|^2, \end{aligned}$$

where $\mathbb{E}_t = \mathbb{E}[\cdot \mid \mathcal{F}_t]$ denotes the conditional expectation.

The following is a one-step improvement guarantee for the stochastic gradient method in the two conceptually distinct cases $C_t = 0$ and $B = 0$. In the case $C_t = 0$, the bias $\mathbb{E}_t g^t - G(x^t)$ shrinks as the iterates approach x^* . The theorem shows that as long as $B/\alpha < 1$, with a sufficiently small stepsize η , the method can converge to an arbitrarily small neighborhood of x^* . In the case $B = 0$, one can only hope to convergence to a neighborhood of the minimizer whose radius depends on $\{C_t\}_{t \geq 0}$.

Theorem 24 (One step improvement) *The following are true.*

- **(Benign bias)** *Suppose $C_t \equiv 0$ for all t . Set $\rho := B/\alpha$ and suppose that we are in the regime $\rho \in (0, 1)$. Then with any $\eta < \frac{\alpha(1-\rho)}{8L^2}$, the stochastic gradient method (37) generates a sequence x^t satisfying*

$$\mathbb{E}_t \|x^{t+1} - x^*\|^2 \leq \frac{1 + 2\eta\alpha\rho + 4\eta^2 D^2 + 2\eta^2 \alpha^2 \rho^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})} \|x^t - x^*\|^2 + \frac{4\eta^2 \sigma_t^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})}. \quad (38)$$

- **(Offset bias)** *Suppose $B \equiv 0$. Then with any $\eta \leq \frac{\alpha}{4L^2}$, the stochastic gradient method (37) generates a sequence x^t satisfying*

$$\mathbb{E}_t \|x^{t+1} - x^*\|^2 \leq \frac{1 + 2\eta^2 D^2}{1 + \eta\alpha} \|x^t - x^*\|^2 + \frac{2\eta^2 \sigma_t^2}{1 + \eta\alpha} + \frac{2\eta C_t^2}{\alpha(1 + \eta\alpha)}. \quad (39)$$

Moreover, in the zero bias setting $B \equiv C_t \equiv 0$, the estimate (39) holds in the slightly wider parameter regime $\eta \leq \frac{\alpha}{2L^2}$.

Proof We begin with the first claim. To this end, suppose $C_t \equiv 0$ for all t . Set $\rho := B/\alpha$ and suppose that we are in the regime $\rho \in (0, 1)$. Fix three constants $\Delta_1, \Delta_2, \Delta_3 > 0$ to be specified later. Noting that x^{t+1} is the minimizer of the 1-strongly convex function $x \mapsto \frac{1}{2}\|x^t - \eta g^t - x\|^2$ over $\mathcal{X}$, we deduce

$$\frac{1}{2}\|x^{t+1} - x^\star\|^2 \leq \frac{1}{2}\|x^t - \eta g^t - x^\star\|^2 - \frac{1}{2}\|x^t - \eta g^t - x^{t+1}\|^2.$$

Expanding the squares on the right hand side and combining terms yields

$$\begin{aligned} \frac{1}{2}\|x^{t+1} - x^\star\|^2 &\leq \frac{1}{2}\|x^t - x^\star\|^2 - \eta \langle g^t, x^{t+1} - x^\star \rangle - \frac{1}{2}\|x^{t+1} - x^t\|^2 \\ &= \frac{1}{2}\|x^t - x^\star\|^2 - \eta \langle g^t, x^t - x^\star \rangle - \frac{1}{2}\|x^{t+1} - x^t\|^2 - \eta \langle g^t, x^{t+1} - x^t \rangle. \end{aligned}$$

Setting $\mu^t := \mathbb{E}_t[g^t]$, we successively compute

$$\begin{aligned} \frac{1}{2}\mathbb{E}_t\|x^{t+1} - x^\star\|^2 &\leq \frac{1}{2}\|x^t - x^\star\|^2 - \eta \langle \mathbb{E}_t g^t, x^t - x^\star \rangle - \frac{1}{2}\mathbb{E}_t\|x^{t+1} - x^t\|^2 - \eta \mathbb{E}_t \langle g^t, x^{t+1} - x^t \rangle \\ &\leq \frac{1}{2}\|x^t - x^\star\|^2 - \eta \langle \mu^t, x^t - x^\star \rangle - \frac{1}{2}\mathbb{E}_t\|x^{t+1} - x^t\|^2 - \eta \mathbb{E}_t \langle g^t, x^{t+1} - x^t \rangle \\ &= \frac{1}{2}\|x^t - x^\star\|^2 - \eta \mathbb{E}_t \langle G(x^{t+1}), x^{t+1} - x^\star \rangle - \frac{1}{2}\mathbb{E}_t\|x^{t+1} - x^t\|^2 \\ &\quad + \underbrace{\eta \mathbb{E}_t \langle g^t - \mu^t, x^t - x^{t+1} \rangle}_{P_1} + \underbrace{\eta \mathbb{E}_t \langle \mu^t - G(x^{t+1}), x^\star - x^{t+1} \rangle}_{P_2}. \end{aligned} \quad (40)$$

Taking into account strong monotonicity of G , we deduce $\langle G(x^{t+1}), x^{t+1} - x^\star \rangle \geq \alpha \|x^{t+1} - x^\star\|^2$ and therefore

$$\frac{1 + 2\eta\alpha}{2}\mathbb{E}_t\|x^{t+1} - x^\star\|^2 \leq \frac{1}{2}\|x^t - x^\star\|^2 - \frac{1}{2}\mathbb{E}_t\|x^{t+1} - x^t\|^2 + \eta(P_1 + P_2). \quad (41)$$

Using Young's inequality, we may upper bound P_1 and P_2 by:

$$P_1 \leq \frac{\sigma_t^2 + D^2\|x^t - x^\star\|^2}{2\Delta_1} + \frac{\Delta_1\mathbb{E}_t\|x^{t+1} - x^t\|^2}{2}. \quad (42)$$

Next, we decompose P_2 as

$$P_2 = \langle \mu^t - G(x^t), x^\star - x^t \rangle + \mathbb{E}_t \langle \mu^t - G(x^t), x^t - x^{t+1} \rangle + \mathbb{E}_t \langle G(x^t) - G(x^{t+1}), x^\star - x^{t+1} \rangle. \quad (43)$$

We bound each of the three products in turn. The first bound follows from our assumption on the bias:

$$\langle \mu^t - G(x^t), x^\star - x^t \rangle \leq B\|x^t - x^\star\|^2. \quad (44)$$

The second bound uses Young's inequality and the assumption on the bias:

$$\begin{aligned} \mathbb{E}_t \langle \mu^t - G(x^t), x^t - x^{t+1} \rangle &\leq \frac{\Delta_2\|\mu^t - G(x^t)\|^2}{2} + \frac{\mathbb{E}_t\|x^t - x^{t+1}\|^2}{2\Delta_2} \\ &\leq \frac{\Delta_2 B^2\|x^t - x^\star\|^2}{2} + \frac{\mathbb{E}_t\|x^t - x^{t+1}\|^2}{2\Delta_2} \end{aligned} \quad (45)$$

The third inequality uses Young's inequality and Lipschitz continuity of G :

$$\begin{aligned}\mathbb{E}_t \langle G(x^t) - G(x^{t+1}), x^* - x^{t+1} \rangle &\leq \frac{\Delta_3 \|G(x^t) - G(x^{t+1})\|^2}{2} + \frac{\mathbb{E}_t \|x^* - x^{t+1}\|^2}{2\Delta_3} \\ &\leq \frac{\Delta_3 L^2 \|x^t - x^{t+1}\|^2}{2} + \frac{\mathbb{E}_t \|x^* - x^{t+1}\|^2}{2\Delta_3}\end{aligned}\quad (46)$$

Putting together all the estimates (41)-(46) yields

$$\begin{aligned}\frac{1 + 2\eta\alpha - 2\eta\Delta_3^{-1}}{2} \mathbb{E}_t \|x^{t+1} - x^*\|^2 &\leq \frac{1 + \eta D^2 \Delta_1^{-1} + 2\eta B + \eta \Delta_2 B^2}{2} \|x^t - x^*\|^2 \\ &\quad - \frac{1 - \eta \Delta_1 - \eta \Delta_2^{-1} - \eta \Delta_3 L^2}{2} \mathbb{E}_t \|x^{t+1} - x^t\|^2 + \frac{\eta \sigma_t^2}{2\Delta_1}.\end{aligned}\quad (47)$$

Let us now set

$$\Delta_3^{-1} = \frac{(1 - \rho)\alpha}{2}, \quad \Delta_1 = \frac{1}{4\eta}, \quad \Delta_2^{-1} := \eta^{-1} - \Delta_1 - \Delta_3 L^2.$$

Notice $\Delta_1 \leq \frac{1}{2\eta} - \Delta_3 L^2$ by our assumption that $\eta \leq \frac{\alpha(1-\rho)}{8L^2}$. In particular, this implies $\Delta_2^{-1} \geq \frac{1}{2\eta}$. Notice that our choice of Δ_2 ensures that the the fraction multiplying $\mathbb{E}_t \|x^{t+1} - x^t\|^2$ in (47) is zero. We therefore deduce

$$\begin{aligned}\mathbb{E}_t \|x^{t+1} - x^*\|^2 &\leq \frac{1 + \eta D^2 \Delta_1^{-1} + 2\eta B + \eta \Delta_2 B^2}{1 + 2\eta\alpha - 2\eta\Delta_3^{-1}} \|x^t - x^*\|^2 + \frac{\eta \sigma_t^2}{\Delta_1(1 + 2\eta\alpha - 2\eta\Delta_3^{-1})} \\ &\leq \frac{1 + 2\eta\alpha\rho + 4\eta^2 D^2 + 2\eta^2 \alpha^2 \rho^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})} \|x^t - x^*\|^2 + \frac{4\eta^2 \sigma_t^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})},\end{aligned}$$

thereby completing the proof of (38).

We next prove the second claim. To this end, suppose $B = 0$. All of the reasoning leading up to and including (42) is valid. Continuing from this point, using Young's inequality, we upper bound P_2 by:

$$P_2 \leq \frac{\mathbb{E}_t \|\mu^t - G(x^{t+1})\|^2}{2\Delta_2} + \frac{\Delta_2 \mathbb{E}_t \|x^{t+1} - x^*\|^2}{2}.\quad (48)$$

Next observe

$$\begin{aligned}\mathbb{E}_t \|\mu^t - G(x^{t+1})\|^2 &\leq 2\mathbb{E}_t \|\mu^t - G(x^t)\|^2 + 2\mathbb{E}_t \|G(x^t) - G(x^{t+1})\|^2, \\ &\leq 2C_t^2 + 2L^2 \|x^t - x^{t+1}\|^2\end{aligned}\quad (49)$$

and therefore

$$P_2 \leq \frac{2C_t^2 + 2L^2 \|x^t - x^{t+1}\|^2}{2\Delta_2} + \frac{\Delta_2 \mathbb{E}_t \|x^{t+1} - x^*\|^2}{2}\quad (50)$$

Putting the estimates (41), (42), and (50) together yields:

$$\begin{aligned}\frac{1 + \eta(2\alpha - \Delta_2)}{2} \mathbb{E}_t \|x^{t+1} - x^*\|^2 &\leq \frac{1 + \eta D^2 \Delta_1^{-1}}{2} \|x^t - x^*\|^2 \\ &\quad + \frac{\eta \sigma_t^2}{2\Delta_1} + \frac{2\eta C_t^2 \Delta_2^{-1}}{2} - \frac{1 - 2\eta L^2 \Delta_2^{-1} - \eta \Delta_1}{2} \mathbb{E}_t \|x^{t+1} - x^t\|^2\end{aligned}\quad (51)$$

Setting $\Delta_2 = \alpha$ and $\Delta_1 = \eta^{-1} - \frac{2L^2}{\alpha}$ ensures that the last term on the right is zero. Notice that our assumption $\eta \leq \frac{\alpha}{4L^2}$ ensures $\Delta_1 \geq \frac{1}{2\eta}$. Rearranging (51) directly yields (39). In the case $B \equiv C_t \equiv 0$, instead of (49) we may simply use the bound $\mathbb{E}_t \|\mu^t - G(x^{t+1})\|^2 = \mathbb{E}_t \|G(x^t) - G(x^{t+1})\|^2 \leq L^2 \|x^t - x^{t+1}\|^2$. Continuing in the same way as above yields the improved estimate. $\blacksquare$

Appendix B. Technical results on sequences

The following lemma is standard and follows from a simple inductive argument.

Lemma 25 *Consider a sequence $D_t \geq 0$ for $t \geq 1$ and constants $t_0 \geq 0$, $a > 0$ satisfying*

$$D_{t+1} \leq (1 - \frac{2}{t+t_0})D_t + \frac{a}{(t+t_0)^2} \quad (52)$$

Then the estimate holds:

$$D_t \leq \frac{\max\{(1+t_0)D_1, a\}}{t+t_0} \quad \forall t \geq 1. \quad (53)$$

We will need the following more general version of the lemma.

Lemma 26 *Consider a sequence $D_t \geq 0$ for $t \geq 1$ and constants $t_0 \geq 0$, $a, b > 0$ satisfying*

$$D_{t+1} \leq (1 - \frac{2}{t+t_0})D_t + \frac{a}{(t+t_0)^2} + \frac{b}{(t+t_0)^3}. \quad (54)$$

Then the estimate holds:

$$D_t \leq \frac{\max\{(1+t_0)D_1/2, a\}}{t+t_0} + \frac{\max\{(1+t_0)^{3/2}D_1/2, b\}}{(t+t_0)^{3/2}} \quad \forall t \geq 1. \quad (55)$$

Proof Clearly the recursion (54) continues to hold with a and b replaced by the bigger quantities $\max\{(1+t_0)D_1/2, a\}$ and $\max\{(1+t_0)D_1/2, b\}$, respectively. Therefore abusing notation, we will make this substitution. As a consequence, the claimed estimate (55) holds automatically for the base case $t = 1$. As an inductive assumption, suppose the claim (55) is true for D_t . Set $s = t + t_0$. We then deduce

$$\begin{aligned} D_{t+1} &\leq \left(1 - \frac{2}{s}\right) D_t + \frac{a}{s^2} + \frac{b}{s^3} \\ &\leq \left(1 - \frac{2}{s}\right) \left(\frac{a}{s} + \frac{b}{s^{3/2}}\right) + \frac{a}{s^2} + \frac{b}{s^3} \\ &\leq a \left(\frac{1}{s} - \frac{1}{s^2}\right) + b \left(\frac{1}{s^{3/2}} - \frac{2}{s^{5/2}} + \frac{1}{s^3}\right). \end{aligned}$$

Elementary algebraic manipulations show $\frac{1}{s} - \frac{1}{s^2} \leq \frac{1}{s+1}$. Define the function $\phi(s) = \frac{1}{s^{3/2}} - \frac{2}{s^{5/2}} + \frac{1}{s^3} - \frac{1}{(1+s)^{3/2}}$. Elementary calculus shows that ϕ is increasing on the interval $s \in [1, \infty)$. Since ϕ tends to zero as s tends to infinity, it follows that ϕ is negative on the interval $[1, \infty)$. We therefore conclude $D_{t+1} \leq \frac{a}{s+1} + \frac{b}{(1+s)^{3/2}}$ as claimed. $\blacksquare$

Appendix C. Online Least Squares

In this appendix section, we record basic and well-known results on estimation in online least squares, following (Dieuleveut et al., 2017).

Lemma 27 *Fix a probability space $(\Omega, \mathcal{F}, \mathbb{P})$ with two sub- σ -algebras $\mathcal{G}_1 \subset \mathcal{G}_2 \subset \mathcal{F}$. Define the function*

$$f(B) = \frac{1}{2} \|By - b\|^2,$$

where $B: \Omega \rightarrow \mathbb{R}^{m_1 \times m_2}$, $b: \Omega \rightarrow \mathbb{R}^{m_1}$, and $y: \Omega \rightarrow \mathbb{R}^{m_2}$ are random variables. Suppose moreover that there exist random variables $V: \Omega \rightarrow \mathbb{R}^{m_1 \times m_2}$ and $\sigma: \Omega \rightarrow \mathbb{R}$ satisfying the following.

1. B , V , and σ_1 are $\mathcal{G}_1$ -measurable.
2. y is $\mathcal{G}_2$ -measurable.
3. The estimates, $\mathbb{E}[b \mid \mathcal{G}_2] = Vy$ and $\mathbb{E}[\|Vy - b\|^2 \mid \mathcal{G}_2] \leq \sigma^2$, hold.
4. There exist constants $\lambda_1, \lambda_2, R > 0$ satisfying

$$\lambda_1 I \preceq \mathbb{E}[yy^\top \mid \mathcal{G}_1], \quad \mathbb{E}[\|y\|^2 \mid \mathcal{G}_1] \leq \lambda_2, \quad \text{and} \quad \mathbb{E}[\|y\|^2 yy^\top \mid \mathcal{G}_1] \preceq R^2 \mathbb{E}[yy^\top].$$

Then for any constant $\nu \in (0, \frac{2}{R^2})$, the gradient step $B^+ = B - \nu(By - b)y^\top$ satisfies the bound:

$$\frac{1}{2} \mathbb{E}[\|B^+ - V\|_F^2 \mid \mathcal{G}_1] \leq \frac{1 - \lambda_1 \nu (2 - \nu R^2)}{2} \|B - V\|_F^2 + \frac{\nu^2 \sigma^2 \lambda_2}{2}.$$

Proof Expanding the squared norm yields:

$$\begin{aligned} \frac{1}{2} \|B^+ - V\|_F^2 &= \frac{1}{2} \|B - V - \nu(By - b)y^\top\|_F^2 = \frac{1}{2} \|B - V\|_F^2 - \nu \langle B - V, (By - b)y^\top \rangle \\ &\quad + \frac{\nu^2}{2} \|(By - b)y^\top\|_F^2. \end{aligned}$$

Taking conditional expectations, we conclude

$$\begin{aligned} \frac{1}{2} \mathbb{E}[\|B^+ - V\|_F^2 \mid \mathcal{G}_2] &= \frac{1}{2} \|B - V\|_F^2 - \nu \langle B - V, (By - \mathbb{E}[b \mid \mathcal{G}_2])y^\top \rangle + \frac{\nu^2}{2} \mathbb{E}[\|(By - b)y^\top\|_F^2 \mid \mathcal{G}_2] \\ &= \frac{1}{2} \|B - V\|_F^2 - \nu \|(B - V)y\|_F^2 + \frac{\nu^2}{2} \|y\|^2 \mathbb{E}[\|By - b\|_F^2 \mid \mathcal{G}_2]. \end{aligned} \tag{56}$$

Next, observe

$$\|By - b\|_F^2 = \|(B - V)y\|^2 + \|Vy - b\|^2 + 2\langle By - Vy, Vy - b \rangle.$$

Taking the conditional expectation $\mathbb{E}[\cdot \mid \mathcal{G}_2]$, the last term vanishes, and therefore we deduce $\mathbb{E}[\|By - b\|_F^2 \mid \mathcal{F}'] \leq \|(B - V)y\|^2 + \sigma^2$. Combining this with (56) we compute

$$\frac{1}{2} \mathbb{E}[\|B^+ - V\|_F^2 \mid \mathcal{G}_2] \leq \frac{1}{2} \|B - V\|_F^2 - \nu \|(B - V)y\|_F^2 + \frac{\nu^2}{2} \|y\|^2 \|(B - V)y\|^2 + \frac{\nu^2 \sigma^2}{2} \|y\|^2.$$

Taking expectations with respect to $\mathcal{G}_1$ and using the tower rule, we deduce

$$\begin{aligned} \frac{1}{2}\mathbb{E}[\|B^+ - V\|_F^2 \mid \mathcal{G}_1] &\leq \frac{1}{2}\|B - V\|_F^2 - \nu \mathbb{E}[\|(B - V)y\|_F^2 \mid \mathcal{G}_1] + \frac{\nu^2}{2} \mathbb{E}[\|y\|^2 \|(B - V)y\|^2 \mid \mathcal{G}_1] \\ &\quad + \frac{\nu^2 \sigma^2 \lambda_2}{2}. \end{aligned}$$

Observe next

$$\mathbb{E}[\|y\|^2 \|(B - V)y\|^2 \mid \mathcal{G}_1] = \langle (B - V)(B - V)^\top, \mathbb{E}[\|y\|^2 yy^\top \mid \mathcal{G}_1] \rangle \leq R^2 \mathbb{E}[\|(B - V)y\|_F^2 \mid \mathcal{G}_1],$$

and therefore

$$\frac{1}{2}\mathbb{E}[\|B^+ - V\|_F^2 \mid \mathcal{G}_1] \leq \frac{1}{2}\|B - V\|_F^2 - \left(\nu - \frac{\nu^2 R^2}{2}\right) \mathbb{E}[\|(B - V)y\|_F^2 \mid \mathcal{G}_1] + \frac{\nu^2 \sigma^2 \lambda_2}{2}$$

Note that $\nu \geq \frac{\nu^2 R^2}{2}$. Next we estimate

$$\mathbb{E}[\|(B - V)y\|_F^2 \mid \mathcal{G}_1] = \text{tr}((B - V)^\top (B - V) \mathbb{E}[yy^\top \mid \mathcal{G}_1]) \geq \lambda_1 \|B - V\|_F^2.$$

This completes the proof. ■

Appendix D. Efficiency of Different Equilibrium Outcomes

In practice, players may not all be accounting for performativity (i.e., running Nash seeking algorithms from Section 6). It is more likely the case that some players are behaving myopically. For example, some players may not take any performativity into account (i.e., run a repeated retraining algorithm such as repeated stochastic gradient method from Section 4), or they may only take into account their own performative effects and not those of their competitors in which case they may be running a stochastic gradient method that accounts only for decision dependence as a function of x_i —i.e., such players model $\mathcal{D}_i(x_i)$ as opposed to $\mathcal{D}_i(x_i, x_{-i})$.

Recall from Section 3 that the gradient of $\mathcal{L}_i(x_i, x_{-i}) = \mathbb{E}_{z_i \sim \mathcal{D}_i(x_i, x_{-i})} \ell_i(x_i, z_i)$ is

$$\nabla_i \mathcal{L}_i(x_i, x_{-i}) = \underbrace{\mathbb{E}_{z_i \sim \mathcal{D}_i(x)} [\nabla_i \ell_i(x_i, x_{-i}, z_i)]}_{G_{i,x}} + \underbrace{\nabla_{u_i} \left(\mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x_{-i})} [\ell_i(x_i, x_{-i}, z_i)] \right)}_{H_{i,x}} \Big|_{u_i=x_i}, \quad (57)$$

When player i is running repeated stochastic gradient descent they are only using an estimate of $G_{i,x}$. On the other hand, when a player is running stochastic gradient descent on $\mathcal{L}_i$ they are using and estimate of $G_{i,x}(x) + H_{i,x}(x)$.

In each of these cases it is the $H_{i,x}$ term which is changing depending on the decision model of the player: either they ignore it completely, or take it into account fully. We comment on the partial information case (partially myopic) below.

Mixed Myopic-Strategic Dynamics. Consider an n -player game where n_2 players are myopic and employing a repeated stochastic gradient method (cf. §4), and the other $n_1 = n - n_2$ players are strategic, accounting for performative effects and using a Nash seeking algorithm (cf. §6). With appropriately chosen step sizes the combined dynamics will converge to a stable point. Indeed this follows from Theorem 24 under the assumption that the mapping

$$\tilde{G}_x(x) = G_x(x) + \underbrace{(H_{1,x}(x), \dots, H_{n_1,x}(x), 0, \dots, 0)}_{\tilde{H}_x(x)}$$

is strongly monotone and the appropriate statistical assumptions hold for individual players as required by the method of individual gradient play they are employing.¹²

Let $\tilde{x}$ be the solution to the following variational inequality:

$$\langle -\tilde{G}_{\tilde{x}}(\tilde{x}), x - \tilde{x} \rangle \leq 0 \quad \forall x \in \mathcal{X}. \quad (58)$$

That is $\tilde{x}$ is an asymptotically stable attractor of the mixed myopic-strategic dynamics described above.

D.1 Bounding the Error to Performatively Stable and Nash Equilibrium

We first bound the distance of the point $\tilde{x}$ to the other equilibrium concepts.

Lemma 28 (Deviation between stable and Nash equilibria) *Suppose that Assumptions 1 and 2 hold and that we are in the regime $\rho < 1$. Moreover, suppose that the expression (57) is valid and the loss functions $\ell_i(\cdot, x_{-i}, z_i)$ are L_i -Lipschitz continuous on $\mathcal{X}_i$. Consider an n -player game such that n_2 players are myopic and $n_1 = n - n_2$ players are strategic as described above. Let $\tilde{x}$ solve (58), and let x^{ps} and x^{ne} be, respectively, a performatively stable equilibrium and a Nash equilibrium. Then the following estimates hold:*

$$\|\tilde{x} - x^{\text{ne}}\| \leq \frac{\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^{n_1} L_i^2 \gamma_i^2}}{\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}} \quad \text{and} \quad \|\tilde{x} - x^{\text{ps}}\| \leq \frac{\sqrt{\sum_{i=1}^{n_1} L_i^2 \gamma_i^2}}{\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}.$$

Proof We tackle each claimed bound separately.

Bound on deviation from Nash equilibrium. Using strong monotonicity of $G_x(x)$ and the fact that $\tilde{x}$ solves (58), we compute

$$\begin{aligned} \alpha \|\tilde{x} - x^{\text{ne}}\|^2 &\leq \langle G_{\tilde{x}}(x^{\text{ne}}) - G_{\tilde{x}}(\tilde{x}), x^{\text{ne}} - \tilde{x} \rangle, \\ &\leq \langle G_{\tilde{x}}(x^{\text{ne}}), x^{\text{ne}} - \tilde{x} \rangle + \langle \tilde{H}_{\tilde{x}}(\tilde{x}), x^{\text{ne}} - \tilde{x} \rangle, \\ &\leq \langle G_{\tilde{x}}(x^{\text{ne}}) - G_{x^{\text{ne}}}(x^{\text{ne}}) + G_{x^{\text{ne}}}(x^{\text{ne}}), x^{\text{ne}} - \tilde{x} \rangle + \sqrt{\sum_{i=1}^{n_1} L_i^2 \gamma_i^2} \cdot \|x^{\text{ne}} - \tilde{x}\|, \\ &\leq \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2} \cdot \|x^{\text{ne}} - \tilde{x}\|^2 + \langle G_{x^{\text{ne}}}(x^{\text{ne}}), x^{\text{ne}} - \tilde{x} \rangle + \sqrt{\sum_{i=1}^{n_1} L_i^2 \gamma_i^2} \cdot \|x^{\text{ne}} - \tilde{x}\|, \end{aligned} \quad (59)$$

12. We do not repeat all of these cases here; the reader can look back to Sections 4–6.

where the last inequality follows from Lemma 5. Next recall that by definition of Nash equilibrium and the expression (6) we have

$$0 \in G_{x^{\text{ne}}}(x^{\text{ne}}) + H_{x^{\text{ne}}}(x^{\text{ne}}) + N_{\mathcal{X}}(x^{\text{ne}}). \quad (60)$$

Note that the Kantorovich-Rubenstein dual representation for W_1 distance directly implies $\|H_{x^{\text{ne}}}(x^{\text{ne}})\| \leq \sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2}$. Therefore we deduce that

$$\langle G_{x^{\text{ne}}}(x^{\text{ne}}), x^{\text{ne}} - \tilde{x} \rangle \leq \sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} \cdot \|x^{\text{ne}} - \tilde{x}\|.$$

Combining (59)-(60) and rearranging and dividing by $\|x^{\text{ne}} - \tilde{x}\|$ gives the claimed bound on $\|x^{\text{ne}} - \tilde{x}\|$.

Bound on deviation from performatively stable equilibrium. Using strong monotonicity of $G_x(x)$ and the fact that $\tilde{x}$ solves (58), we compute

$$\begin{aligned} \alpha \|\tilde{x} - x^{\text{ps}}\|^2 &\leq \langle G_{\tilde{x}}(x^{\text{ps}}) - G_{\tilde{x}}(\tilde{x}), x^{\text{ps}} - \tilde{x} \rangle, \\ &\leq \langle G_{\tilde{x}}(x^{\text{ps}}), x^{\text{ps}} - \tilde{x} \rangle + \langle \tilde{H}_{\tilde{x}}(\tilde{x}), x^{\text{ps}} - \tilde{x} \rangle, \\ &\leq \langle G_{\tilde{x}}(x^{\text{ps}}) - G_{x^{\text{ps}}}(x^{\text{ne}}) + G_{x^{\text{ps}}}(x^{\text{ps}}), x^{\text{ne}} - \tilde{x} \rangle + \sqrt{\sum_{i=1}^{n_1} L_i^2 \gamma_i^2} \cdot \|x^{\text{ps}} - \tilde{x}\|, \\ &= \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2} \cdot \|x^{\text{ps}} - \tilde{x}\|^2 + \sqrt{\sum_{i=1}^{n_1} L_i^2 \gamma_i^2} \cdot \|x^{\text{ne}} - \tilde{x}\|. \end{aligned}$$

where we used the fact that $\langle -G_{x^{\text{ps}}}(x^{\text{ps}}), \tilde{x} - x^{\text{ps}} \rangle \leq 0$. Rearranging and dividing by $\|x^{\text{ps}} - \tilde{x}\|$ yields the claimed bound. $\blacksquare$

Note that players may have partial information about Q_i . For instance, consider a player is running stochastic gradient descent only accounting for their own performative effects they are using an estimate of

$$\mathbb{E}_{z_i \sim \mathcal{D}_i(x)} [\nabla_i \ell_i(x_i, x_{-i}, z_i)] + \nabla_{u_i} \left(\mathbb{E}_{z_i \sim \tilde{\mathcal{D}}_i(u_i)} [\ell_i(x_i, x_{-i}, z_i)] \right) \Big|_{u_i = x_i}, \quad (61)$$

where notice that this player is using an estimate $\tilde{\mathcal{D}}_i(u_i)$ of $\mathcal{D}_i(x_i, x_{-i})$. For example, if $\mathcal{D}_i$ is a location-scale distribution such that $z_i = \zeta_i + A_i x_i + A_{-i} x_{-i}$, then $\tilde{\mathcal{D}}_i$ may be modeled as $\tilde{z}_i = \zeta_i + A_i x_i$. Analogous results can be stated for this case as well under some assumptions on $\tilde{\mathcal{D}}_i$ —i.e., that it satisfies Assumption 5. For the sake of brevity we do not include this case.

D.2 Bounding the Distance to the Social Optimum

Under further regularity conditions we can also bound the distance between the social optimum and the different equilibrium. In particular, we need that the losses ℓ_i are not just

L_i -Lipschitz continuous on $\mathcal{X}_i$ but on all of $\mathcal{X}$. First, let us define the social optimum. The social cost (amongst the players, not including any abstracted utility of the users represented by the decision dependent distribution) is the sum of the losses of all the players—i.e.,

$$\mathcal{C}(x) = \sum_{i=1}^n \mathcal{L}_i(x_i, x_{-i})$$

and a social optimum is a minimizer of $\mathcal{C}(x)$. When this loss is strongly convex, the social optimum is unique. Let

$$x^* \in \operatorname{argmin}_{x \in \mathcal{X}} \mathcal{C}(x).$$

In particular x^* solves the variational inequality

$$\langle -S_{x^*}(x^*), x - x^* \rangle \leq 0 \quad \forall x \in \mathcal{X}, \quad (62)$$

where $S_x(x) = G_x(x) + H_x(x) + G'_x(x) + H'_x(x)$ such that

$$\begin{aligned} G'_x(x) &= \left(\mathbb{E}_{z_i \sim \mathcal{D}_i(x)} \nabla_{-i} \ell_i(x, z_i) \right)_{i=1}^n \\ H'_x(x) &= \left(\nabla_{u_{-i}} \left(\mathbb{E}_{z_i \sim \mathcal{D}_i(x_i, u_{-i})} [\ell_i(x_i, x_{-i}, z_i)] \Big|_{u_j = x_j} \right) \right) \end{aligned}$$

Lemma 29 (Distance to social optimum) *Suppose that Assumptions 1, 2, and 5 hold and that we are in the regime $\rho < 1/2$. Moreover, suppose that the expression (57) is valid and the loss functions $\ell_i(\cdot, x_{-i}, z_i)$ are L_i -Lipschitz continuous on $\mathcal{X}$. Let x^{so} solve (62), and let x^{ps} and x^{ne} be, respectively, a performatively stable equilibrium and a Nash equilibrium. Then the following estimates hold:*

$$\|x^{\text{so}} - x^{\text{ne}}\| \leq \frac{\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^n L_i^2}}{(1-2\rho)\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}, \quad \text{and} \quad \|x^{\text{so}} - x^{\text{ps}}\| \leq \frac{2\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^n L_i^2}}{(1-2\rho)\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}.$$

Proof We show the proof for the distance between the Nash equilibrium and the social optimum and note that the proof for bounding the distance between x^{ps} and x^* is analogous.

From Theorem 13 that the game $\mathcal{G}$ (defined in (5)) is $(1-2\rho)\alpha$ strongly monotone. Let $\bar{G}_x(x) := G_x(x) + H_x(x)$. Using strong monotonicity of $\bar{G}_x(x)$ and the fact that x^{so} solves (62), we compute

$$\begin{aligned} (1-2\rho)\alpha \|x^{\text{so}} - x^{\text{ne}}\|^2 &\leq \langle \bar{G}_{x^{\text{ne}}}(x^{\text{so}}) - \bar{G}_{x^{\text{ne}}}(x^{\text{ne}}), x^{\text{so}} - x^{\text{ne}} \rangle \\ &\leq \langle \bar{G}_{x^{\text{ne}}}(x^{\text{so}}) - \bar{G}_{x^{\text{so}}}(x^{\text{so}}) + \bar{G}_{x^{\text{so}}}(x^{\text{so}}), x^{\text{so}} - x^{\text{ne}} \rangle \\ &\leq \langle G_{x^{\text{ne}}}(x^{\text{so}}) - G_{x^{\text{so}}}(x^{\text{so}}) + H_{x^{\text{ne}}}(x^{\text{so}}) - H_{x^{\text{so}}}(x^{\text{so}}), x^{\text{so}} - x^{\text{ne}} \rangle \\ &\quad + \langle \bar{G}_{x^{\text{so}}}(x^{\text{so}}), x^{\text{so}} - x^{\text{ne}} \rangle \\ &\leq \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2} \cdot \|x^{\text{so}} - x^{\text{ne}}\|^2 + \langle \bar{G}_{x^{\text{so}}}(x^{\text{so}}), x^{\text{so}} - x^{\text{ne}} \rangle, \end{aligned}$$

■

where the last inequality holds since $H_x(y)$ is strongly monotone in x . Now since x^{so} satisfies (62), we have that

$$(1 - 2\rho)\alpha \|x^{\text{so}} - x^{\text{ne}}\|^2 \leq \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2} \cdot \|x^{\text{so}} - x^{\text{ne}}\|^2 + \left(\sqrt{\sum_{i=1}^n L_i^2} + \sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} \right) \|x^{\text{so}} - x^{\text{ne}}\|$$

so that rearranging and dividing through by $\|x^{\text{so}} - x^{\text{ne}}\|$, the claimed bound holds.

Finally we can also bound the distance between the social optimum and the myopic-strategic equilibrium $\tilde{x}$ defined in the preceding subsection.

Lemma 30 (Distance to social optimum) *Suppose that Assumptions 1, 2, and 5 hold and that we are in the regime $\rho < 1/2$. Moreover, suppose that the expression (57) is valid and the loss functions $\ell_i(\cdot, x_{-i}, z_i)$ are L_i -Lipschitz continuous on $\mathcal{X}$. Consider an n -player game such that n_2 players are myopic and $n_1 = n - n_2$ players are strategic as described above. Let x^{so} solve (62), and let $\tilde{x}$ solve (58). Then the following estimate holds:*

$$\|x^{\text{so}} - \tilde{x}\| \leq \frac{2\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^n L_i^2}}{(1 - 2\rho)\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}.$$

D.3 Bounding Player Losses

We can theoretically bound how much worse off a player is in terms of their loss using the lemmas in the proceeding section and Lemma 28.

Corollary 31 *Suppose that Assumptions 1 and 2 hold and that we are in the regime $\rho < 1$. Moreover, suppose that the expression (57) is valid and the loss functions $\ell_i(\cdot, x_{-i}, z_i)$ are L_i -Lipschitz continuous on $\mathcal{X}_i$. Let $\tilde{x}$ solve (58) (in the setting where n_2 players are myopic and the rest are strategic), and let x^{ps} , x^{ne} , and x^{so} be, respectively, a performatively stable equilibrium, a Nash equilibrium, and a social optimum. Suppose player i plays x_i^{ps} while the others play is consistent with x_{-i}^{ne} . Then the following hold:*

- A player myopically playing any strategy x'_i other than x_i^{ne} including x_i^{so} , $\tilde{x}_i$, and x_i^{ps} is worse off than in the Nash equilibrium assuming all other players play is consistent with x^{ne} .*
- If $x'_i \in \{x_i^{\text{so}}, x_i^{\text{ps}}, \tilde{x}_i\}$, the difference between the experienced loss and the loss at the Nash equilibrium for all other (strategic) players is bounded.*
- If $x'_i \in \{x_i^{\text{ne}}, x_i^{\text{ps}}, \tilde{x}_i\}$, the difference between the experienced loss and the loss at the social optimum for all other (strategic) players is bounded.*

Proof To see that a. holds recall the definition of Nash. Indeed, by the definition of a Nash equilibrium, if player i plays x'_i and all other players play x_{-i}^{ne} , then

$$\mathcal{L}_i(x^{\text{ne}}) \leq \mathcal{L}_i(x'_i, x_{-i}^{\text{ne}}).$$

That is, they are worse off. To see that b. holds, for the strategic players $j \neq i$ we have that

$$|\mathcal{L}_j(x^{\text{ne}}) - \mathcal{L}_j(x_{-i}^{\text{ne}}, x'_i)| \leq L_j \|x'_i - x_{-i}^{\text{ne}}\| \leq L_j \cdot c_j,$$

where

$$c_j = \begin{cases} \frac{\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2}}{\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}, & x'_i = x_i^{\text{ps}} \\ \frac{\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^{n-1} L_i^2 \gamma_i^2}}{\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}, & x'_i = \tilde{x}_i \\ \frac{\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^n L_i^2}}{(1-2\rho)\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}, & x'_i = x_i^{\text{so}}. \end{cases} \quad (63)$$

To see that c. holds, for the strategic players $j \neq i$ we have that

$$|\mathcal{L}_j(x^{\text{so}}) - \mathcal{L}_j(x_{-i}^{\text{so}}, x'_i)| \leq L_j \|x'_i - x_{-i}^{\text{so}}\| \leq L_j \cdot c'_j,$$

where

$$c'_j = \begin{cases} \frac{2\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^n L_i^2}}{(1-2\rho)\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}, & x'_i = x_i^{\text{ps}} \\ \frac{2\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^n L_i^2}}{(1-2\rho)\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}, & x'_i = \tilde{x}_i \\ \frac{\sqrt{\sum_{i=1}^n L_i^2 \gamma_i^2} + \sqrt{\sum_{i=1}^n L_i^2}}{(1-2\rho)\alpha - \sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}}, & x'_i = x_i^{\text{ne}}. \end{cases} \quad (64)$$

■

Appendix E. Numerical Examples

E.1 Revenue Maximization: Competition in Ride-Share Markets

In this section, we explore is semi-synthetic competition between two ride-share platforms seeking to maximize their revenue given that the demand they experience is influenced by their own prices as well as their competitors. We use data from a prior Kaggle competition to set up the semi-synthetic simulation environment.¹³

E.2 Game Abstraction

We repeat the game abstraction for the ride-share example here for ease of access. The abstraction for the game can be described as follows. Consider a ride-share market with two platforms that each seek to maximize their revenue by setting prices for their rides given by a vector $x_i \in \mathbb{R}^{m_i}$. The vector of demands $z_i \in \mathbb{R}^{m_i}$ containing demand information for m_i locations that each ride-share platform serves is influenced not only by the prices they set but also the prices that their competitor sets. Suppose that platform i 's loss is given by

$$\ell_i(x_i, z_i) = -\frac{1}{2} z_i^\top x_i + \frac{\lambda_i}{2} \|x_i\|^2$$

where $\lambda_i \geq 0$ is some regularization parameter, and $x_i \in \mathbb{R}^{m_i}$ represents the vector of price differentials from some nominal price for each of the m_i locations. Observe that this game is

13. The data used in this paper is publicly available (<https://www.kaggle.com/br11rb/uber-and-lyft-dataset-boston-ma>).

separable since the losses ℓ_i do not explicitly depend on x_{-i} . Moreover, let us suppose that the random demand z_i takes the semi-parametric form $z_i = \zeta_i + A_i x_i + A_{-i} x_{-i}$, where ζ_i follows some base distribution $\mathcal{P}_i$ and the parameters A_i and A_{-i} capture price elasticities to player i 's and its competitor's change in price, respectively; naturally, the price elasticity for player i to its own price changes is negative while the price elasticity for player i 's demand given changes in its competitors actions is positive. Namely, we have that $A_i \preceq 0$ and $A_{-i} \succeq 0$ capturing that an increase in player i 's prices results in a decrease in demand, while an increase in its competitor's prices results in a increase in demand. Moreover, we showed in Example 1 that the mapping $x \mapsto H_x(y)$ is trivially monotone. Hence, the game between ride-share platforms is strongly monotone and admits a unique Nash equilibrium. Throughout the remainder of this section we set $\lambda_1 = \lambda_2 = 1$.

E.3 Semi-Synthetic Data Construction

There are eleven locations that we consider in our simulation, and each element in x_i represents the price difference (set by platform i) from a nominal price at each location. We aggregate the rides into bins of \$5 increments; this is done by taking the raw data and rounding the price to the nearest bin as follows: $p_{\text{nominal}} = 5 \cdot \lfloor \frac{p}{5} \rfloor$ where p is the actual price of a particular ride. Then, for each bin j we have an empirical distribution $\mathcal{P}_{i,j}$ for each player $i \in \{1, 2\}$ which is just the collection of rides in the data set for the location and price range specified by that bin.

Ride-Share Simulation Parameter Estimation. In the experiments presented, we estimate the matrices A_i and A_{-i} that govern the performative effects from the data. We estimate the price elasticity matrices A_i for each player from the data using the heuristic that an increase in price by a factor of a by either firm leads to a decrease by a factor of b in their own demand. In order to achieve this behavior, the diagonal elements of A_i are set as follows in terms of the price for the bin and the expected base demand taken from the empirical distribution for the bin. That is,

$$[A_i]_{jj} \cdot a \cdot p_j = (1 - b)\bar{\xi}_j \implies [A_i]_{jj} = \frac{(1 - b)\bar{\xi}_j}{a \cdot p_j}, \quad (65)$$

where p_j is the nominal price assigned to bin j and $\bar{\xi}_j$ is the expected value of the base demand sampled from the empirical distribution of rides in bin j .

We construct A_{-i} using a similar heuristic and the empirical average demands for each bin. In the experiments, we started with the assumption that a 50% increase in prices by one firm would decrease that firm's demand by 75% while increasing the competitor's demand by 37.5%. This can be interpreted as an assumption that when one firm increases prices they will lose some customers, half of whom will switch services, and half of whom will not use either service. Therefore, for the matrix A_{-i} , the diagonal elements are half the size of the diagonal elements of A_i . In the linked code base, we provide a mechanism to set the off-diagonals to simulate correlation between locations, however, we do not include simulations in the paper due to length.

In the code base it is possible to vary the values of a and b in (65) to achieve more or less competition between the players; when the values on the diagonal of A_{-i} are large relative to the elements on the diagonal of A_i , the effects of the players on each other are

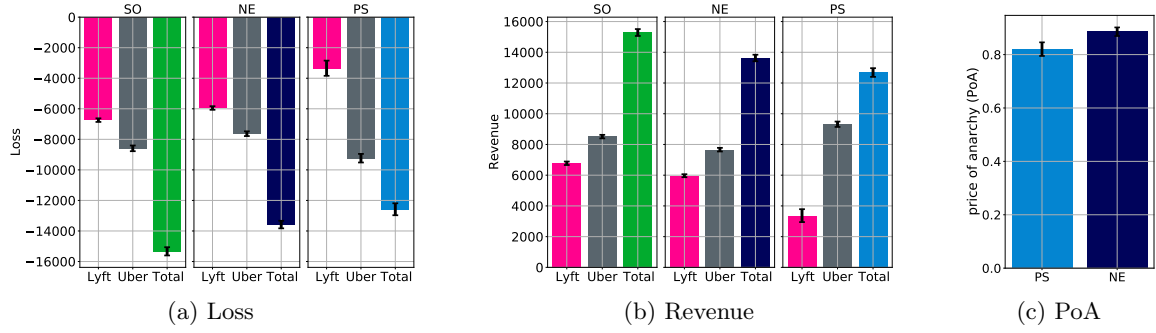


Figure 4: **Competition in Ride-Share Markets: Experiment 2.** (a) Loss and (b) revenue at the social optimum obtained via stochastic gradient descent on the social cost (sum of players’ costs), Nash equilibrium obtained via the stochastic gradient method, and performatively stable equilibrium obtained via the repeated stochastic gradient method. The overall (sum of both players) loss and revenue are worse at the performatively stable equilibrium. Note that the loss is the negative of the revenue plus some small ($\lambda_1 = \lambda_2 = 1$) regularization term. (c) Average price of anarchy (PoA) at the Nash equilibrium versus the performatively stable equilibrium. A value closer to one is better, and hence the Nash equilibrium (by a small margin) has a better PoA.



Figure 5: **Competition in Ride-Share Markets: Experiment 3.** Effects due to myopic decision-making. Change in (a) price, (b) demand and (c) revenue from nominal by location for \$10 nominal price bin due to ignoring performative effects: players run stochastic gradient descent, and the image shows the change in demand (respectively, revenue) when both players model decision-dependence as compared to when they both do not model decision-dependence. The size of the circles shows the magnitude of the change, while the color indicates the raw value. The majority of locations see a *decrease* in demand, but due to an increase in price at the Nash equilibrium relative to the myopic outcome, there is an increase in revenue for *both* players.

larger than their effects on themselves, so the competition dominates. On the other hand, if the elements of A_{-i} are zero or negligible compared to A_i , the game reduces to separate performative prediction problems for the two players.

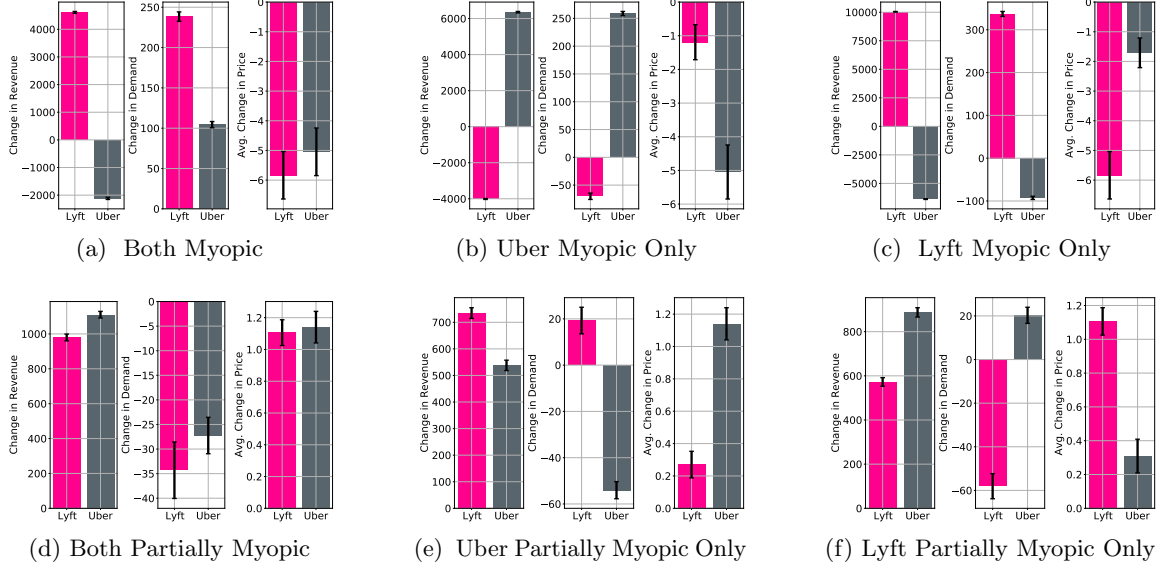


Figure 6: **Competition in Ride-Share Markets: Experiment 3.** Effects of players being (a)–(c) myopic or (d)–(f) partially myopic relative to Nash (not myopic, and consider competition). Positive changes in revenue indicate the Nash equilibrium is better for that player. When a player is myopic, they do not consider any performative effects in their update—i.e., $g_i^t = \lambda_i x_i^t - \frac{1}{2} \zeta_i^t$ —and when a player is partially myopic, they consider their own performative effects, but not those of their competitor—i.e., $g_i^t = -(A_i - \lambda_i I)^\top x_i^t - \frac{1}{2} \zeta_i^t$. In (a)–(c), we observe that when at least one player is completely myopic, then at least one player is worse off at the Nash equilibrium. In (d)–(f) we observe that when at least one player is partially myopic, the Nash equilibrium always is better for both players.

Experiment D.1: Effect of Ignoring Performativity. We study the impact of players ignoring performative effects due to competition in the data distribution. In Figure 6, we explore the effects of players either being completely myopic—i.e., $g_i^t = \lambda_i x_i^t - \frac{1}{2} \zeta_i^t$ —or partially myopic—i.e., $g_i^t = -(A_i - \lambda_i I)^\top x_i^t - \frac{1}{2} \zeta_i^t$ —on the change in revenue, demand and average price (across locations) from nominal at the Nash equilibrium. Recall that players employing the stochastic gradient method use the gradient estimate $g_i^t = -(A_i - \lambda_i I)^\top x_i^t - \frac{1}{2} (\zeta_i^t + A_{-i} x_{-i}^t)$; we refer to this as the non-myopic case since all performative effects are considered. Even when the players are myopic or partially myopic, the environment, however, does have these performative effects, meaning that $z_i = \zeta_i + A_i x_i + A_{-i} x_{-i}$ and hence, the myopic player is in this sense ignoring or unaware of the fact that the data distribution is reacting to its competition’s decisions.

In Figure 6 (a)–(c), we observe that when at least one player is completely myopic, then at least one player is worse off at the Nash equilibrium in the sense that their revenue is lower. This is theoretical justified by Corollary 31.a. In Figure 6 (d)–(f), on the other hand, we observe that when at least one player is partially myopic, the Nash equilibrium always is better for both players in the sense that their individual revenues are higher at the Nash.

The values in Figure 6 represent the total demand and revenue changes, and average price change across locations. It is also informative to examine the per-location changes. Focusing in on the setting considered in Figure 6 (d), we examine the per-location price, revenue and demand. We see that the relative change depends on the location, however, the majority of locations see a decrease in demand, yet an increase in price and hence, revenue. This suggests that modeling performative effects due to competition can be beneficial for both players.

References

- S. Ahmed. *Strategic planning under uncertainty: Stochastic integer programming approaches*. University of Illinois at Urbana-Champaign, 2000.
- J. Angwin, J. Larson, S. Mattu, and L. Kirchner. Machine bias. *Propublica*, 2016. URL <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>.
- R. Bartlett, A. Morse, R. Stanton, and N. Wallace. Consumer-lending discrimination in the fintech era. Technical report, National Bureau of Economic Research, 2019.
- Y. Bechavod, K. Ligett, Z. S. Wu, and J. Ziani. Causal feature discovery through strategic modification. *arXiv preprint arXiv:2002.07024*, 3, 2020.
- M. Bravo, D. Leslie, and P. Mertikopoulos. Bandit learning in concave n-person games. *Advances in Neural Information Processing Systems*, 31, 2018.
- G. Brown, S. Hod, and I. Kalemaj. Performative prediction in a stateful world. In *International Conference on Artificial Intelligence and Statistics*, pages 6045–6061. PMLR, 2022.
- B. Chasnov, L. Ratliff, E. Mazumdar, and S. Burden. Convergence analysis of gradient-based learning in continuous games. In R. P. Adams and V. Gogate, editors, *Proceedings of The 35th Uncertainty in Artificial Intelligence Conference*, volume 115 of *Proceedings of Machine Learning Research*, pages 935–944. PMLR, 22–25 Jul 2020a.
- B. Chasnov, L. J. Ratliff, E. Mazumdar, and S. A. Burden. Convergence analysis of gradient-based learning with non-uniform learning rates in non-cooperative multi-agent settings. *Proceedings of the Conference on Uncertainty in Artificial Intelligence (UAI)*, 2020b.
- L. Chen, A. Mislove, and C. Wilson. Peeking beneath the hood of uber. In *Proceedings of the 2015 Internet Measurement Conference*, IMC ’15, page 495–508, 2015. doi: 10.1145/2815675.2815681.
- R. Courtland. Bias detectives: the researchers striving to make algorithms fair. *Nature*, pages 357–360, 2018. doi: 10.1038/d41586-018-05469-3.
- J. Cutler, D. Drusvyatskiy, and Z. Harchaoui. Stochastic optimization under time drift: iterate averaging, step-decay schedules, and high probability guarantees. *Advances in neural information processing systems*, 34:11859–11869, 2021.

- S. Dean, M. Curmei, L. J. Ratliff, J. Morgenstern, and M. Fazel. Multi-learner risk reduction under endogenous participation dynamics. *arXiv preprint arXiv:2206.02667*, 2022.
- T. Diethe, T. Borchert, E. Thereska, B. Balle, and N. Lawrence. Continual learning in practice. *Continual Learning Workshop of 32nd Conference on Neural Information Processing Systems (NeurIPS 2018)*, 2019.
- A. Dieuleveut, N. Flammarion, and F. Bach. Harder, better, faster, stronger convergence rates for least-squares regression. *The Journal of Machine Learning Research*, 18(1): 3520–3570, 2017.
- J. Dong, A. Roth, Z. Schutzman, B. Waggoner, and Z. S. Wu. Strategic classification from revealed preferences. In *Proceedings of the 2018 ACM Conference on Economics and Computation*, pages 55–70, 2018.
- D. Drusvyatskiy and L. Xiao. Stochastic optimization with decision-dependent distributions. *Mathematics of Operations Research*, 48(2):954–998, 2023.
- D. Drusvyatskiy, M. Fazel, and L. J. Ratliff. Improved rates for derivative free gradient play in strongly monotone games. In *2022 IEEE 61st Conference on Decision and Control (CDC)*, pages 3403–3408. IEEE, 2022.
- J. Dupacová. Optimization under exogenous and endogenous uncertainty. *University of West Bohemia in Pilsen*, 2006.
- D. Ensign, S. A. Friedler, S. Neville, C. Scheidegger, and S. Venkatasubramanian. Runaway feedback loops in predictive policing. In *Conference on fairness, accountability and transparency*, pages 160–171. PMLR, 2018.
- Y. Ermoliev. Stochastic quasigradient methods. In Y. Ermoliev and R. J.-B. Wets, editors, *Numerical Techniques for Stochastic Optimization*, chapter 6, pages 141–185. Springer, 1988.
- A. A. Gaivoronskii. Nonstationary stochastic programming problems. *Cybernetics*, 14: 575–579, 1978.
- S. Ghadimi and G. Lan. Optimal stochastic approximation algorithms for strongly convex stochastic composite optimization, ii: shrinking procedures and optimal algorithms. *SIAM Journal on Optimization*, 23(4):2061–2089, 2013.
- M. Hardt, N. Megiddo, C. Papadimitriou, and M. Wootters. Strategic classification. In *Proceedings of the 2016 ACM conference on innovations in theoretical computer science*, pages 111–122, 2016.
- L. Hellemo, P. I. Barton, and A. Tomasgard. Decision-dependent probabilities in stochastic programs with recourse. *Computational Management Science*, 15(3):369–395, 2018.
- Z. Huo and H. Huang. Asynchronous mini-batch gradient descent with variance reduction for non-convex optimization. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 31 (1), 2017.

- Z. Izzo, L. Ying, and J. Zou. How to learn when data reacts to your model: performative gradient descent. *Proceedings of the conference on Artificial Intelligence and Statistics (AISTats)*, 2022.
- T. W. Jonsbråten, R. J. Wets, and D. L. Woodruff. A class of stochastic programs with decision dependent random elements. *Annals of Operations Research*, 82:83–106, 1998.
- Q. Li and H.-T. Wai. State dependent performative prediction with stochastic approximation. In *International Conference on Artificial Intelligence and Statistics*, pages 3164–3186. PMLR, 2022.
- X. Lian, Y. Huang, Y. Li, and J. Liu. Asynchronous parallel stochastic gradient for nonconvex optimization. *Advances in neural information processing systems*, 28, 2015.
- K. Lum and W. Isaac. To predict and serve? *Significance*, 13(5):14–19, 2016.
- C. Mendler-Dünnér, J. Perdomo, T. Zrnic, and M. Hardt. Stochastic optimization for performative prediction. In H. Larochelle, M. Ranzato, R. Hadsell, M. F. Balcan, and H. Lin, editors, *Advances in Neural Information Processing Systems*, volume 33, pages 4929–4939. Curran Associates, Inc., 2020.
- P. Mertikopoulos and Z. Zhou. Learning in games with continuous action sets and unknown payoff functions. *Mathematical Programming*, 173(1):465–507, 2019.
- J. Miller, S. Milli, and M. Hardt. Strategic classification is causal modeling in disguise. In *International Conference on Machine Learning*, pages 6917–6926. PMLR, 2020.
- J. Miller, J. C. Perdomo, and T. Zrnic. Outside the echo chamber: Optimizing the performative risk. *arXiv preprint arXiv:2102.08570*, 2021.
- J. Perdomo, T. Zrnic, C. Mendler-Dünnér, and M. Hardt. Performative prediction. In *International Conference on Machine Learning*, pages 7599–7609. PMLR, 2020.
- L. J. Ratliff, S. A. Burden, and S. S. Sastry. On the characterization of local nash equilibria in continuous games. *IEEE transactions on automatic control*, 61(8):2301–2307, 2016.
- M. Ray, L. J. Ratliff, D. Drusvyatskiy, and M. Fazel. Decision-dependent learning in geometrically decaying environments. In *Proceedings of the AAAI International Conference on Artificial Intelligence (AAAI)*, 2022.
- B. Recht, C. Re, S. Wright, and F. Niu. Hogwild!: A lock-free approach to parallelizing stochastic gradient descent. *Advances in neural information processing systems*, 24, 2011.
- J. B. Rosen. Existence and uniqueness of equilibrium points for concave n-person games. *Econometrica: Journal of the Econometric Society*, pages 520–534, 1965.
- R. Y. Rubinstein and A. Shapiro. *Discrete event systems: Sensitivity analysis and stochastic optimization by the score function method*, volume 13. Wiley, 1993.
- T. Tatarenko and M. Kamgarpour. Learning nash equilibria in monotone games. In *2019 IEEE 58th Conference on Decision and Control (CDC)*, pages 3104–3109. IEEE, 2019.

- T. Tatarenko and M. Kamgarpour. Bandit online learning of nash equilibria in monotone games. *arXiv preprint arXiv:2009.04258*, 2020.
- P. Varaiya and R.-B. Wets. Stochastic dynamic optimization approaches and computation. *International Institute for Applied Systems Analysis, Working paper WP-88-087*, 1988.
- H. Von Stackelberg. *Market structure and equilibrium*. Springer Science & Business Media, 2010.
- K. Wood, G. Bianchin, and E. Dall’Anese. Online projected gradient descent for stochastic optimization with decision-dependent distributions. *IEEE Control Systems Letters*, 2021.
- Y. Wu, E. Dobriban, and S. Davidson. Deltagrad: Rapid retraining of machine learning models. In *International Conference on Machine Learning*, pages 10355–10366. PMLR, 2020.
- Z. Zhou, P. Mertikopoulos, S. Athey, N. Bambos, P. Glynn, and Y. Ye. Learning in games with lossy feedback. In *Proceedings of the 32nd Conference on Neural Information Processing Systems (NeurIPS)*, pages 1–11, 2018.