
Learning in Stochastic Monotone Games with Decision-Dependent Data

Adhyayan Narang Evan Faulkner Dmitriy Drusvyatskiy Maryam Fazel Lillian J. Ratliff
University of Washington

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 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.

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 since the underlying data distribution may “shift” in reaction to the decision maker’s actions. Indeed, machine learning algorithms are increasingly being trained on data that is generated by strategic or even adversarial agents, and then deployed in environments that react in response 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 shift is to

periodically retrain the model to adapt to the changing distribution (Diethe et al., 2019; Wu et al., 2020). Despite the ubiquity of retraining heuristics in practice, 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 literature on strategic classification (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 a growing body of research that develops algorithms for performative prediction by leveraging advances in convex optimization (Drusvyatskiy and Xiao, 2020; Miller et al., 2021; Mendler-Dünnér et al., 2020; Perdomo et al., 2020; Brown et al., 2020). A more extensive literature review is contained in Appendix A.

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. On the other hand, learning algorithms in practice are often deployed alongside other algorithms which may even be competing with one another. One concrete example to keep in mind is that of university admissions, wherein applicants may tailor their profile to make them more desirable for the college of their choice. In this case, there are multiple competing decision-makers (universities) and the population of applicants reacts based on the admissions policies of all the universities simultaneously. Examples of this type are widespread in applications and we provide more detailed vignettes in Section 3.

Contributions. We formulate the first game theoretic model for decision-dependent learning in the presence of competition, which we call *multi-player performative prediction*. This is a new class of stochastic games with relevance to many machine learning tasks.

Aiming towards algorithms for finding Nash equilibria for multiplayer performative prediction games, we develop transparent conditions ensuring strong monotonicity of the game. Assuming that the game is indeed strongly monotone, we discuss a number of algorithms for finding Nash equilibria. In particular, derivative-free methods are immediately applicable but have a high sample complexity $\mathcal{O}(\frac{d^2}{\varepsilon})$. 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 directly applicable and comes equipped with an efficiency guarantee of $\mathcal{O}(\frac{d}{\varepsilon})$.

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, in each iteration, the players simultaneously estimate the parameters of their decision-dependent distribution and optimize their loss by taking a step in the direction of their individual gradient, 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 $\mathcal{O}(\frac{d}{\varepsilon})$ of the stochastic gradient method.

Finally, we present illustrative numerical experiments using a semi-synthetic example that uses data from multiple ride-share companies (Section 6).

2 NOTATION & 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 $\mathcal{X} \subset \mathbb{R}^d$ will be denoted by $\text{proj}_{\mathcal{X}}(y) = \text{argmin}_{x \in \mathcal{X}} \|x - y\|$. The normal cone to a convex set $\mathcal{X}$ at $x \in \mathcal{X}$ is the set $N_{\mathcal{X}}(x) = \{v \in \mathbb{R}^d : \langle v, y - x \rangle \leq 0 \ \forall y \in \mathcal{X}\}$.

Convex Games and Monotonicity. Fix an index set $[n] = \{1, \dots, n\}$, integers d_i for $i \in [n]$, and set

$d = \sum_{i=1}^n d_i$. We will always 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

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

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}^* .

We use the symbol $\nabla_i \mathcal{L}(x)$ to denote the derivative of $\mathcal{L}(\cdot)$ with respect to x_i . With this notation, we define the vector of individual gradients

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

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, and the partial gradients $\nabla_i \mathcal{L}_i(x)$ exist and are continuous. Thus, the Nash equilibria x^* are characterized by the inclusion

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

Generally speaking, finding global Nash equilibria is only possible for “monotone” games. A C^1 -smooth convex game is *α -strongly monotone* (for $\alpha \geq 0$) if

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

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

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

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:

$$W_1(\mu, \nu) = \sup_{h \in \text{Lip}_1} \left\{ \mathbb{E}_{X \sim \mu} [h(X)] - \mathbb{E}_{Y \sim \nu} [h(Y)] \right\}, \quad (2)$$

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

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

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 (2020, Lemmas 1.1, 2.1) or (Perdomo et al., 2020, Lemma C.4) 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 -smooth for all $z \in \mathcal{Z}$ and the map $z \mapsto \nabla_x g(x, z)$ is β -Lipschitz continuous for any $x \in \mathbb{R}^d$. Then for any measures $\mu, \nu \in \mathbb{P}(\mathcal{Z})$, 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

$$\min_{x_i \in \mathcal{X}_i} \mathcal{L}_i(x_i, x_{-i}) \text{ where } \mathcal{L}_i(x) := \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} \ell_i(x, z_i). \quad (3)$$

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 $\mathcal{D}_i(x) \in \mathbb{P}(\mathcal{Z}_i)$ 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. Appendix A). We leave it to future work to investigate both layers of strategic interaction simultaneously.

We assume that each player has full information of the other players’ parameters. 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. 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. The decision x_i player i makes is the rule for recommending routes. Users choose routes, which then impact the realized travel time $z \sim \mathcal{D}$ on the road segments in the network observed by all players.

Vignette 2 (Multiplayer Strategic Classification). Players have different distributions—i.e., $\mathcal{D}_i \neq \mathcal{D}_j$ for all $i, j \in [n]$. Multiple universities classify students as accepted or rejected using applicant data, where each applicant designs their application to fit desiderata of the different 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. 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 .

Prior formulations of decision dependent learning do not readily extend to the settings described in the vignettes without a game theoretic model. The classical notion of Nash equilibrium is a natural equilibrium concept for the game (3).

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

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

Generally speaking, finding Nash equilibria is only computationally feasible for monotone games. The game (3) can easily fail to be monotone even if the game is separable and the loss functions $\ell_i(\cdot, z)$ are strongly convex. In Section 4, we develop sufficient conditions for (strong) monotonicity and use them to analyze algorithms for finding Nash equilibria.

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

Assumption 1 (Convexity and smoothness). *There exist constants $\alpha > 0$ and $\beta_i > 0$ such that for each $i \in [n]$, the following hold: (i) For any $y \in \mathcal{X}$, the game $\mathcal{G}(y)$ is α -strongly monotone. (ii) Each loss $\ell_i(x_i, x_{-i}, z_i)$ is C^1 -smooth in x_i and the map $z_i \mapsto \nabla_{x_i} \ell_i(x_i, x_{-i}, z_i)$ is β_i -Lipschitz continuous for any $x \in \mathcal{X}$.*

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 (3) itself. Sufficient conditions for convexity and strong monotonicity of the game are given in Section 4.

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*

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

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

We end this section with some convenient notation that will be used throughout.

Notation. To this end, 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 $g_i(x, z_i) := \nabla_{x_i} \ell_i(x, z_i)$ and $g(x, z) := (g_1(x, z_1), \dots, g_n(x, z_n))$. Taking expectations define $G_{i,y}(x) := \mathbb{E}_{z_i \sim \mathcal{D}_i(y)} g_i(x, z_i)$ and

$$G_y(x) := (G_{1,y}(x), \dots, G_{n,y}(x)). \quad (4)$$

We may also express G_y as $G_y(x) := \mathbb{E}_{z \sim \mathcal{D}_\pi(y)} g(x, z)$ where $\mathcal{D}_\pi(y) := \mathcal{D}_1(y) \times \dots \times \mathcal{D}_n(y)$ is the product measure. The following lemma bounding the deviation in the vector of individual gradients is a direct consequence of Lemma 1.

Lemma 4. *Suppose Assumptions 1 and 2 hold. For every $x, y, y' \in \mathcal{X}$ and $i \in [n]$, the estimates hold:*

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

4 MONOTONICITY

Towards developing algorithms for finding true Nash equilibria of the game (3), 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, and necessarily so because the game (3) 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 3 (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 3, the chain rule implies the derivative of player i 's loss function with respect to their own choice variable x_i is given by $\nabla_{x_i} \mathcal{L}_i(x_i, x_{-i}) = G_{i,x}(x) + H_{i,x}(x)$, where

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

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 corresponding to the game (3) is simply the map $D(x) := G_x(x) + H_x(x)$. Thus the game (3) is monotone, as long as $D(x)$ is a monotone mapping.

The sufficient conditions we present in Theorem 6 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 (i.e., $n = 1$), 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). The proof of Theorem 6 crucially relies on the following.

Lemma 5. *Suppose that Assumptions 1, 2, 3 hold. For any $x \in \mathcal{X}$, the map $H_x(y)$ is Lipschitz continuous in y with parameter $\sqrt{\sum_{i=1}^n \beta_i^2 \gamma_i^2}$.*

Theorem 6 (Monotonicity of the decision-dependent game). *Suppose that Assumptions 1–3 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 (3) is strongly monotone with parameter $(1 - 2\rho)\alpha$.*

Theorem 6 gives sufficient conditions under which the game (3) is strongly monotone and hence, admits a unique Nash equilibrium.

The following example of a multiplayer performative prediction problem illustrates settings where the mapping $x \mapsto H_x(y)$ is monotone and therefore Theorem 6 may be applied.

Example 1 (Multiplayer Revenue Maximization). Consider a setting with two firms that each would like to maximize their revenue by setting the price x_i (e.g., a ride-share market). The demand z_i that each 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 $H_{i,x}(y) = -A_i^\top y_i$. Hence, the map $x \mapsto H_x(y)$ is constant and is therefore trivially monotone.

5 ALGORITHMS

In this section we analyze algorithms that converge to the Nash equilibrium of the n -player performative prediction game (3) when the game is strongly monotone. Recall that the Nash equilibrium x^* of this game is characterized by the relation

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

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—seeking Nash equilibrium of continuous games in different information settings. Specifically, we study gradient-based learning methods where players update their strategies using an estimate of their individual gradient consistent with the information available to them.

The Nash-seeking algorithms studied in this section all use gradient estimates of the individual gradient $\nabla_i \mathcal{L}_i(x_i, x_{-i}) = G_{i,x}(x) + H_{i,x}(x)$ for each player $i \in [n]$. The main difficulty with applying gradient-based methods is estimation of the term $H_{i,x}(x)$, without some parametric assumptions on the distributions $\mathcal{D}_i$. Consequently, we start in Section 5.1 with derivative free methods; here, each player only has access to loss function queries, which implies that players do not require access to $\mathcal{D}_i$. To improve efficiency, we then study stochastic gradient methods with different assumptions on oracle access to $\mathcal{D}_i$.

5.1 Derivative Free Method

The first information setting we consider is the “bandit feedback” setting: 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 fewest 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, and such methods have been extended to games (Bravo et al., 2018; Drusvyatskiy et al., 2021). The result in this section is a direct consequence 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 update:

$$\left\{ \begin{array}{l} \text{Sample } v_i^t \in \mathbb{S}_i \\ \text{Sample } z_i^t \sim \mathcal{D}_i(x^t + \delta v^t) \\ \text{Set } x_i^{t+1} = \operatorname{proj}_{(1-\delta)\mathcal{X}_i}(x_i^t - \eta_t g_i^t) \end{array} \right\} \quad (5)$$

where $g_i^t := \frac{d_i}{\delta} \ell_i(x^t + \delta v^t, z_i^t) v_i^t$. 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 $\mathcal{X}_i$. The formal statement for derivative free methods in general games can be found

in Drusvyatskiy et al. (2021).¹

Proposition 7 (Convergence rate of the derivative free method). *Consider an n -player decision-dependent game (3). Under reasonable smoothness and bounded variance assumptions, algorithm (5) 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. In the remainder, we focus on stochastic gradient methods, which enjoy significantly better efficiency guarantees albeit at cost of access to a richer oracle.

5.2 Stochastic Gradient Method

In practice, players often 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.

Assumption 4 (Parametric assumption). *For each index $i \in [n]$, there exists a probability measure $\mathcal{P}_i$ and matrices A_i and A_{-i} satisfying*

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

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

Assumption 4 is very natural and generalizes an analogous assumption used in the single player setting in (Miller et al., 2021). 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 the decisions of player i and all other players $-i$, respectively, on the distribution $\mathcal{D}_i$ governing player i ’s data.

Under Assumption 4, player i ’s loss is given by

$$\mathcal{L}_i(x) = \mathbb{E}_{\zeta_i \sim \mathcal{P}_i} \ell_i(x, \zeta_i + A_i x_i + A_{-i} x_{-i}). \quad (6)$$

¹Though Theorem 2 in Drusvyatskiy et al. (2021) is stated for deterministic games, it applies verbatim whenever the value of the loss function for each player is replaced by an unbiased estimator of their individual loss functions—our setting.

Under mild smoothness assumptions, differentiating (6) using the chain rule, we see that the gradient of the i -th player’s loss is simply

$$\nabla_i \mathcal{L}_i(x) = \mathbb{E}_{z_i \sim \mathcal{D}_i(x)} [\nabla_i \ell_i(x, z_i) + A_i^\top \nabla_{z_i} \ell_i(x, z_i)]. \quad (7)$$

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:

$$\left\{ \begin{array}{l} \text{Sample } z_i^t \sim \mathcal{D}_i(x^t) \\ \text{Set } x_i^{t+1} = \text{proj}_{\mathcal{X}_i}(x_i^t - \eta_t \cdot w_i(x^t, z_i^t)) \end{array} \right\}. \quad (8)$$

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.

Assumption 5 (Smoothness). *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 8 (Sufficient conditions for Assumption 5). *Suppose Assumption 4 holds and for each i there exist constants $\xi_i \geq 0$ such that $(x, z_i) \mapsto \nabla_{i,z_i} \ell_i(x, z_i)$ is ξ_i -Lipschitz continuous. Then, Assumption 5 holds with $L = (\sum_{i=1}^n \xi_i^2 \max\{1, \|A_i\|_{\text{op}}^2\} \cdot (1 + \|A_i\|_{\text{op}}^2))^{1/2}$.*

Next we assume a finite variance bound.

Assumption 6 (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 6 to hold in terms of the variance of the individual gradients $\nabla_{i, z_i} \ell(x, z_i)$. The proof is immediate and we omit it.

Lemma 9 (Sufficient conditions for Assumption 6). *Suppose there exist constants $s_1, s_2 \geq 0$ such that for all $x \in \mathcal{X}$ and $i \in [n]$ the estimates hold:*

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

Then, Assumption 6 holds with $\sigma^2 = \sum_{i=1}^n 2(s_1^2 + \|A_i\|_{\text{op}}^2 s_2^2)$.

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

Theorem 10 (Stochastic gradient play). *Consider an n -player performative prediction game (3). Suppose that Assumptions 4-6 hold and that the game is α -strongly monotone with $\alpha > 0$. Then a single step of the stochastic gradient method (8) with any constant $\eta \leq \frac{\alpha}{2L^2}$ satisfies*

$$\mathbb{E}[\|x^{t+1} - x^*\|^2] \leq \frac{1}{1 + \alpha\eta} \mathbb{E}[\|x^t - x^*\|^2] + \frac{2\eta^2\sigma^2}{1 + \eta\alpha}, \quad (9)$$

where x^* is the Nash equilibrium of the game (3).

This theorem is immediate from Theorem 15 in Appendix E with $B \equiv C_t \equiv D \equiv 0$.

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 (9) and the generic results on step decay schedules; e.g. (Drusvyatskiy and Xiao, 2020, Lemma B.2).

Corollary 11 (Stochastic gradient method with a step-decay schedule). *Suppose that the assumptions of Theorem 10 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*

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

and $K = \lceil 1 + \log_2\left(\frac{2\eta_0\sigma^2}{\alpha\varepsilon}\right) \rceil$. 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 $\mathcal{O}\left(\frac{L^2}{\alpha^2} \cdot \log\left(\frac{2R}{\varepsilon}\right) + \frac{\sigma^2}{\alpha^2\varepsilon}\right)$.*

5.3 Adaptive Gradient Method

Throughout this section, we continue working under the parametric Assumption 4. An apparent deficiency of the stochastic gradient method discussed in Section 5.2 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, where “adaptive” refers to the adaptive estimation scheme employed by players. That is, in each iteration, each player simultaneously estimates their distribution parameters and myopically optimizes their individual loss via the 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) $\hat{A}_i^t$ that dynamically estimates the unknown matrix 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 $x_i^{t+1} = \text{proj}_{\mathcal{X}_i}(x_i^t - \eta_t \hat{g}_i^t)$ where

$$\hat{g}_i^t := \nabla_{i, z_i^t} \ell(x^t, z_i^t) + (\hat{A}_{ii}^t)^\top \nabla_{z_i} \ell(x^t, z_i^t) \quad (10)$$

and $\hat{A}_{ii}^t$ denotes the submatrix of $\hat{A}_i^t$ whose columns are indexed by player i ’s action space. Next, in order to update $\hat{A}_i^t$, 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 $\min_{B_i} \frac{1}{2} \|q_i^t - z_i^t - B_i u^t\|^2$. 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 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. We impose the following assumption.

Assumption 7 (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 $0 \prec c_l \cdot I \preceq \mathbb{E}[v_i v_i^\top]$, $\mathbb{E}\|v_i\|^2 \leq c_{u,i}$, and $\mathbb{E}[\|v_i\|^2 v_i v_i^\top] \preceq R^2 \mathbb{E}[v_i v_i^\top]$.*

In the simple Gaussian case where $u_t \sim \mathcal{N}(0, I_d)$, we may set $c_l = 1$, $c_{u,i} = d_i$, and $R^2 = 3 \max_{i \in [n]} d_i$. An

²For the justification of the expression for R^2 , see

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     Gradient update:
        $x_i^{t+1} = \text{proj}_{\mathcal{X}_i}(x_i^t - \eta_t \hat{g}_i^t)$ ,
       where  $\hat{g}_i^t$  is defined in (10).
6     Estimation update:
        $\hat{A}_i^{t+1} = \hat{A}_i^t + \nu_t(q_i^t - z_i^t - \hat{A}_i^t u_i^t)(u_i^t)^\top$ 
7   end
8 end
9 end
    
```

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 10 in Section 5.2. 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 12 (Estimation error). *Suppose Assumptions 4 and 7 hold and let $\nu_t \in (0, \frac{2}{R^2})$. The matrices $\hat{A}_i^t$ generated by Algorithm 1 satisfy the estimate:*

$$\frac{1}{2} \mathbb{E}_t \|\hat{A}_i^{t+1} - \bar{A}_i\|_F^2 \leq \frac{1 - c_l \nu_t (2 - \nu_t R^2)}{2} \|\hat{A}_i^t - \bar{A}_i\|_F^2 + \nu_t^2 \text{tr}(\Sigma_i) c_{u,i}. \quad (11)$$

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

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

Next we show that the direction of motion of Algorithm 5.3 is well-aligned with the direction of motion of the stochastic gradient method. To this end, define the true (stochastic) vector of individual gradients

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

(Dieuleveut et al., 2017, Section 2.1).

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 8 (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 13. *Suppose Assumptions 4 and 8 hold. For each $t \geq 1$ and $i \in [n]$, the estimate holds:*

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

In light of Lemmas 12 and 13, we may interpret Algorithm 5.3 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 15 of Appendix E. We will make use of the following assumption.

Assumption 9 (Finite variance). *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 14 (Convergence of the adaptive method). *Suppose that Assumptions 4, 5, 7, and 8 hold and that the game (3) 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*

$$\begin{aligned} & \mathbb{E} \|x^t - x^*\|^2 \leq \\ & \frac{\max \left\{ \frac{1}{2} \alpha^2 (1 + k_0) \|x_1 - x^*\|^2, 32 Z_1 \sigma^2 + 8 \delta^2 Z \right\}}{\alpha^2 (t + k_0)} \\ & + \frac{\max \left\{ \frac{1}{2} \alpha^2 (1 + k_0)^{\frac{3}{2}} \|x_1 - x^*\|^2, 64 \sigma^2 Z \right\}}{\alpha^2 (t + k_0)^{\frac{3}{2}}}. \end{aligned}$$

where we set $Z_1 = (1 + 2 \|\bar{A}\|_F^2)$, $Z_2 = \|\hat{A}^1 - \bar{A}\|_F^2$ and

$$Z = \max \left\{ \frac{1+k_0}{1+q_0}, 1 \right\} \cdot \max \left\{ \left(1 + \frac{2R^2}{c_l}\right) Z_2, \frac{8 \sum_{i=1}^n \text{tr}(\Sigma_i) c_{u,i}}{c_l^2} \right\}$$

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

6 SEMISYNTHETIC SIMULATIONS

We consider a semi-synthetic competition between two ride-share platforms.³

Consider a ride-share market with two platforms that each seek to maximize their revenue by setting the price x_i . The vector of demands $z_i \in \mathbb{R}^{m_i}$ containing demand information for m_i locations that each ride-share platform sees 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 $\lambda_i = 1$, $i = 1, 2$ is a regularization parameter, and $x_i \in \mathbb{R}^{m_i}$ represents the vector of price differentials from a 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, we model the random demand z_i in 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. In Appendix H.1 we describe how the data is processed.

To validate the theory developed in the previous sections, we show the iteration complexity of the algorithms in Section 5. We run each algorithm from twenty random initial conditions, and compute the error between the trajectory of the algorithm and the Nash equilibrium. In Figure 1 we show the mean of these error trajectories and plus and minus one standard deviation. For the stochastic gradient method, we use a constant step-size $\eta = 5e-5$, and for the adaptive gradient method we use the step-size schedule $\eta_t = \eta_0/t$ for the gradient update and $\nu_t = \nu_0/t$ for the estimation update where $\eta_0 = 5e-5$ and $\nu_0 = 1e-4$. For the derivative free method, we use a constant query radius $\delta = 5$ and step-size schedule $\eta_t = 2/t$. The plots in Figure 1 demonstrate that empirical iteration complexity of the adaptive gradient method and the

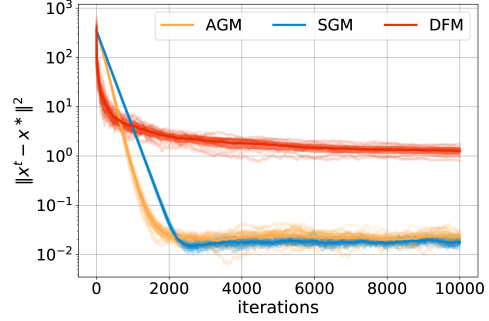


Figure 1: **Error to Nash Equilibrium.** Iteration Complexity for derivative free gradient method (DFM), stochastic gradient method (SGM) and adaptive gradient method (AGM). The plot demonstrates that the iteration complexities of AGM and SGM are nearly identical and outperform DFM as expected.

stochastic gradient method are nearly identical, and outperform the derivative free method as expected.

In Appendix H we provide additional experiments on the effects of ignoring performativity. The general message is that players are better off under the Nash equilibrium arising from modeling performative effects of their competitors.

7 DISCUSSION

The new class of games introduced in this paper motivates interesting future work at the intersection of statistical learning theory 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.

Acknowledgements

Ratliff was supported by NSF CNS-1844729 and Office of Naval Research YIP Award N000142012571. Drusvyatskiy was supported by the NSF DMS 1651851 and CCF 1740551 awards. Fazel was supported by awards NSF TRIPODS II-DMS 2023166, NSF TRIPODS-CCF 1740551, and NSF CCF 2007036.

References

Shabbir Ahmed. *Strategic planning under uncertainty: Stochastic integer programming approaches*. Univer-

³The data can be found linked here.

- sity of Illinois at Urbana-Champaign, 2000.
- Julia Angwin, Jeff Larson, Surya Mattu, and Lauren Kirchner. Machine bias. *ProPublica*, 2016.
- Robert Bartlett, Adair Morse, Richard Stanton, and Nancy Wallace. Consumer-lending discrimination in the fintech era. Technical report, National Bureau of Economic Research, 2019.
- Mario Bravo, David S Leslie, and Panayotis Mertikopoulos. Bandit learning in concave n -person games. *arXiv preprint arXiv:1810.01925*, 2018.
- Gavin Brown, Shlomi Hod, and Iden Kalemaj. Performative prediction in a stateful world. *arXiv preprint arXiv:2011.03885*, 2020.
- Benjamin Chasnov, Lillian Ratliff, Eric Mazumdar, and Samuel Burden. Convergence analysis of gradient-based learning in continuous games. In Ryan P. Adams and Vibhav 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 2020.
- Le Chen, Alan Mislove, and Christo 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.
- Tom Diethe, Tom Borchert, Eno Thereska, Borja Balle, and Neil Lawrence. Continual learning in practice. *arXiv preprint arXiv:1903.05202*, 2019.
- Aymeric Dieuleveut, Nicolas Flammarion, and Francis Bach. Harder, better, faster, stronger convergence rates for least-squares regression. *The Journal of Machine Learning Research*, 18(1):3520–3570, 2017.
- Jinshuo Dong, Aaron Roth, Zachary Schutzman, Bo Waggoner, and Zhiwei Steven Wu. Strategic classification from revealed preferences. In *Proceedings of the 2018 ACM Conference on Economics and Computation*, pages 55–70, 2018.
- Dmitriy Drusvyatskiy and Lin Xiao. Stochastic optimization with decision-dependent distributions. *arXiv preprint arXiv:2011.11173*, 2020.
- Dmitriy Drusvyatskiy, Lillian J. Ratliff, and Maryam Fazel. Improved rates for derivative free gradient play in convex games. *arXiv:2111.09456*, 2021.
- Jitka Dupacová. Optimization under exogenous and endogenous uncertainty. *University of West Bohemia in Pilsen*, 2006.
- Moritz Hardt, Nimrod Megiddo, Christos Papadimitriou, and Mary Wootters. Strategic classification. In *Proceedings of the 2016 ACM conference on innovations in theoretical computer science*, pages 111–122, 2016.
- Lars Hellemo, Paul I Barton, and Asgeir Tomasgard. Decision-dependent probabilities in stochastic programs with recourse. *Computational Management Science*, 15(3):369–395, 2018.
- Zachary Izzo, Lexing Ying, and James Zou. How to learn when data reacts to your model: performative gradient descent. *arXiv preprint arXiv:2102.07698*, 2021.
- Tore W Jonsbråten, Roger JB Wets, and David L Woodruff. A class of stochastic programs with decision dependent random elements. *Annals of Operations Research*, 82:83–106, 1998.
- Kristian Lum and William Isaac. To predict and serve? *Significance*, 13(5):14–19, 2016.
- Celestine Mendler-Dünner, Juan Perdomo, Tijana Zrnic, and Moritz 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.
- Panayotis Mertikopoulos and Zhengyuan Zhou. Learning in games with continuous action sets and unknown payoff functions. *Mathematical Programming*, 173(1):465–507, 2019.
- John Miller, Smitha Milli, and Moritz Hardt. Strategic classification is causal modeling in disguise. In *International Conference on Machine Learning*, pages 6917–6926. PMLR, 2020.
- John Miller, Juan C Perdomo, and Tijana Zrnic. Outside the echo chamber: Optimizing the performative risk. *arXiv preprint arXiv:2102.08570*, 2021.
- Juan Perdomo, Tijana Zrnic, Celestine Mendler-Dünner, and Moritz Hardt. Performative prediction. In *International Conference on Machine Learning*, pages 7599–7609. PMLR, 2020.
- Lillian J Ratliff, Samuel A Burden, and S Shankar Sastry. On the characterization of local nash equilibria in continuous games. *IEEE transactions on automatic control*, 61(8):2301–2307, 2016.
- J Ben Rosen. Existence and uniqueness of equilibrium points for concave n -person games. *Econometrica: Journal of the Econometric Society*, pages 520–534, 1965.
- Reuven Y Rubinstein and Alexander Shapiro. *Discrete event systems: Sensitivity analysis and stochastic optimization by the score function method*, volume 13. Wiley, 1993.

- Tatiana Tatarenko and Maryam Kamgarpour. Learning nash equilibria in monotone games. In *2019 IEEE 58th Conference on Decision and Control (CDC)*, pages 3104–3109. IEEE, 2019.
- Tatiana Tatarenko and Maryam Kamgarpour. Bandit online learning of nash equilibria in monotone games. *arXiv preprint arXiv:2009.04258*, 2020.
- Pravin Varaiya and RJ-B Wets. Stochastic dynamic optimization approaches and computation. *IIASA Working Paper*, 1988.
- Heinrich Von Stackelberg. *Market structure and equilibrium*. Springer Science & Business Media, 2010.
- Yinjun Wu, Edgar Dobriban, and Susan Davidson. Deltagrad: Rapid retraining of machine learning models. In *International Conference on Machine Learning*, pages 10355–10366. PMLR, 2020.

A 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 notions of performative optimality and stability, and show that repeated retraining and stochastic gradient methods converge to a stable point. Subsequently, (Drusvyatskiy and Xiao, 2020) showed that a variety of popular gradient-based algorithms 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, 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. For decision-dependent distributions parameterized by a *location* parameter, (Miller et al., 2021) additionally introduce a two-stage algorithm for finding performatively optimal points. The paper (Izzo et al., 2021) instead focuses on algorithms that estimate gradients with finite differences. It is noteworthy that performative prediction is largely motivated by strategic classification (Hardt et al., 2016).

Gradient-Based Learning in Continuous Games. There is a broad and growing literature on learning in games. We focus here on the most relevant subset: gradient-based learning in continuous games. The seminal work by Rosen (1965) showed that convex games which are *diagonal strictly convex*—i.e., strictly monotone—admit a unique Nash equilibrium and *gradient play*—a gradient method in which each player follows the partial gradient of their cost with respect to their choice variable—converges to it. There is literature extending this work to more general continuous games to obtain a local characterization for equilibria and local convergence guarantees (see, e.g., (Ratliff et al., 2016; Chasnov et al., 2020)). 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., 2021). Relaxing strong monotonicity to monotonicity, by incorporating a regularization term that decays to zero asymptotically, 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.

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.

B Proof for Section 3

Proof of Lemma 4 (Bound on Deviation of Vector of Individual Gradients).

Lemma 4. *Suppose Assumptions 1 and 2 hold. For every $x, y, y' \in \mathcal{X}$ and $i \in [n]$, the estimates hold:*

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

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

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

Therefore, we deduce

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

The proof is complete. $\square$

C Proofs for Section 4

Proof of Lemma 5.

Lemma 5. *Suppose that Assumptions 1, 2, 3 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

$$\|H_{i,x'}(x) - H_{i,x'}(y)\| \leq \int_{s=0}^1 \left\| \frac{d}{du_i} \mathbb{E}_{z_i \sim \mathcal{D}_i(u_i, x'_{-i})} \nabla_i \ell_i(\gamma(s), z_i) \Big|_{u_i=x'_i} \right\| \cdot \|x - y\| ds. \quad (12)$$

Now for any $s \in (0, 1)$, Lemma 4 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 (12) 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 . $\square$

Proof of Theorem 6.

Theorem 6 (Monotonicity of the decision-dependent game). *Suppose that Assumptions 1–3 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 (3) is strongly monotone with parameter $(1 - 2\rho)\alpha$.*

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

$$\langle D(x) - D(x'), x - x' \rangle = \langle G_x(x) - G_{x'}(x'), x - x' \rangle + \langle H_x(x) - H_{x'}(x'), x - x' \rangle. \quad (13)$$

We estimate the first term as follows:

$$\begin{aligned} \langle G_x(x) - G_{x'}(x'), x - x' \rangle &= \langle G_{x'}(x) - G_{x'}(x'), x - x' \rangle + \langle G_x(x) - G_{x'}(x), x - x' \rangle \\ &\geq \alpha \|x - x'\|^2 - \left(\sum_{i=1}^n \beta_i^2 \gamma_i^2 \right)^{1/2} \cdot \|x - x'\|^2 \end{aligned} \quad (14)$$

$$= (1 - \rho) \alpha \|x - x'\|^2, \quad (15)$$

where (14) follows from Assumption 1 and Lemma 4. Next, we estimate the second term on the right side of (13) as follows:

$$\begin{aligned} \langle H_x(x) - H_{x'}(x'), x - x' \rangle &= \langle H_{x'}(x) - H_{x'}(x'), x - x' \rangle + \langle H_x(x) - H_{x'}(x), x - x' \rangle \\ &\geq \langle H_{x'}(x) - H_{x'}(x'), x - x' \rangle \end{aligned} \quad (16)$$

$$\geq -\|H_{x'}(x) - H_{x'}(x')\| \cdot \|x - x'\| \quad (17)$$

$$\geq -\left(\sum_{i=1}^n \beta_i^2 \gamma_i^2 \right)^{1/2} \|x - x'\|^2, \quad (18)$$

where (16) follows from the assumption that the map $x \mapsto H_x(y)$ is monotone and (18) follows from Lemma 5. Combining (13), (15), and (18) completes the proof. $\square$

D Proofs for Section 5

Lemma 8 (Sufficient conditions for Assumption 5.). *Suppose Assumption 4 holds and for each i there exist constants $\xi_i \geq 0$ such that $(x, z_i) \mapsto \nabla_{i,z_i} \ell_i(x, z_i)$ is ξ_i -Lipschitz continuous. Then, Assumption 5 holds with $L = (\sum_{i=1}^n \xi_i^2 \max\{1, \|A_i\|_{\text{op}}^2\} \cdot (1 + \|\bar{A}_i\|_{\text{op}}^2))^{1/2}$.*

Proof of Lemma 8.

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

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

This completes the proof. $\square$

Proof of Lemma 12.

Lemma 12 (Estimation error). *Suppose Assumptions 4 and 7 hold and let $\nu_t \in (0, \frac{2}{R^2})$. The matrices $\hat{A}_i^t$ generated by Algorithm 1 satisfy the estimate:*

$$\begin{aligned} \frac{1}{2} \mathbb{E}_t \|\hat{A}_i^{t+1} - \bar{A}_i\|_F^2 &\leq \frac{1 - c_l \nu_t (2 - \nu_t R^2)}{2} \|\hat{A}_i^t - \bar{A}_i\|_F^2 \\ &\quad + \nu_t^2 \text{tr}(\Sigma_i) c_{u,i}. \end{aligned} \tag{11}$$

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

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

Proof. This follows from a standard estimate for online least squares, which appears as Lemma 18 in Appendix G. 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 18 completes the proof of (11). Summing up (11) 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 16 in Appendix F. $\square$

Proof of Lemma 13.

Lemma 13. *Suppose Assumptions 4 and 8 hold. 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 estimate:

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

as claimed. $\square$

Proof of Theorem 14.

Theorem 14 (Convergence of the adaptive method). *Suppose that Assumptions 4, 5, 7, and 8 hold and that the game (3) 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*

$$\begin{aligned} \mathbb{E} \|x^t - x^*\|^2 &\leq \frac{\max \left\{ \frac{1}{2} \alpha^2 (1 + k_0) \|x_1 - x^*\|^2, 32Z_1 \sigma^2 + 8\delta^2 Z \right\}}{\alpha^2(t + k_0)} \\ &\quad + \frac{\max \left\{ \frac{1}{2} \alpha^2 (1 + k_0)^{\frac{3}{2}} \|x_1 - x^*\|^2, 64\sigma^2 Z \right\}}{\alpha^2(t + k_0)^{\frac{3}{2}}}. \end{aligned}$$

where we set $Z_1 = (1 + 2\|\bar{A}\|_F^2)$, $Z_2 = \|\hat{A}^1 - \bar{A}\|_F^2$ and

$$Z = \max \left\{ \frac{1+k_0}{1+q_0}, 1 \right\} \cdot \max \left\{ \left(1 + \frac{2R^2}{c_l}\right) Z_2, \frac{8 \sum_{i=1}^n \text{tr}(\Sigma_i) c_{u,i}}{c_l^2} \right\}$$

Proof. We will apply the standard convergence guarantees in Theorem 15 of Appendix E for biased stochastic gradient methods. To this end, using Lemma 13 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 15 in Appendix E we deduce

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

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

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

where the last inequality follows from Lemma 12.

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

$$\begin{aligned} \mathbb{E}\|x^{t+1} - x^*\|^2 &\leq \left(1 - \frac{2}{t + k_0}\right) \mathbb{E}\|x^t - x^*\|^2 + \frac{8(1 + 2\|\bar{A}\|_F^2)\sigma^2}{\alpha^2(t + k_0 - 2)^2} + \frac{16\sigma^2 Z}{\alpha^2(t + q_0)(t + k_0 - 2)^2} \\ &\quad + \frac{4\delta^2 Z}{\alpha^2(t + q_0)(t + k_0 - 2)}. \end{aligned} \tag{19}$$

We now aim to apply Lemma 17 in Section F. To this end, we need to upper bound the last three terms in (19) 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 (19), we obtain

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

Applying Lemma 17 in Section F, we conclude:

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

This completes the proof. $\square$

E Stochastic gradient method with bias

In this section, we consider a variational inequality

$$0 \in G(x) + N_{\mathcal{X}}(x), \tag{20}$$

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:

$$x^{t+1} = \text{proj}_{\mathcal{X}}(x^t - \eta g^t), \tag{21}$$

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 (20).

Assumption 10 (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 15 (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 (21) generates a sequence x^t satisfying*

$$\mathbb{E}_t \|x^{t+1} - x^*\|^2 \leq \frac{1 + 2\eta\alpha\rho + 4\eta^2 D^2 + 2\eta^2 \alpha^2 \rho^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})} \|x^t - x^*\|^2 + \frac{4\eta^2 \sigma_t^2}{1 + 2\eta\alpha(\frac{1+\rho}{2})}. \quad (22)$$

- **(Offset bias)** *Suppose $B \equiv 0$. Then with any $\eta \leq \frac{\alpha}{4L^2}$, the stochastic gradient method (21) generates a sequence x^t satisfying*

$$\mathbb{E}_t \|x^{t+1} - x^*\|^2 \leq \frac{1 + 2\eta^2 D^2}{1 + \eta\alpha} \|x^t - x^*\|^2 + \frac{2\eta^2 \sigma_t^2}{1 + \eta\alpha} + \frac{2\eta C_t^2}{\alpha(1 + \eta\alpha)}. \quad (23)$$

Moreover, in the zero bias setting $B \equiv C_t \equiv 0$, the estimate (23) 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

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

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

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

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

$$\begin{aligned} \frac{1}{2}\mathbb{E}_t \|x^{t+1} - x^*\|^2 & \leq \frac{1}{2}\|x^t - x^*\|^2 - \eta \langle \mathbb{E}_t g^t, x^t - x^* \rangle - \frac{1}{2}\mathbb{E}_t \|x^{t+1} - x^t\|^2 - \eta \mathbb{E}_t \langle g^t, x^{t+1} - x^t \rangle \\ & \leq \frac{1}{2}\|x^t - x^*\|^2 - \eta \langle \mu^t, x^t - x^* \rangle - \frac{1}{2}\mathbb{E}_t \|x^{t+1} - x^t\|^2 - \eta \mathbb{E}_t \langle g^t, x^{t+1} - x^t \rangle \\ & = \frac{1}{2}\|x^t - x^*\|^2 - \eta \mathbb{E}_t \langle G(x^{t+1}), x^{t+1} - x^* \rangle - \frac{1}{2}\mathbb{E}_t \|x^{t+1} - x^t\|^2 \\ & \quad + \underbrace{\eta \mathbb{E}_t \langle g^t - \mu^t, x^t - x^{t+1} \rangle}_{P_1} + \underbrace{\eta \mathbb{E}_t \langle \mu^t - G(x^{t+1}), x^* - x^{t+1} \rangle}_{P_2}. \end{aligned} \quad (24)$$

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

$$\frac{1 + 2\eta\alpha}{2} \mathbb{E}_t \|x^{t+1} - x^*\|^2 \leq \frac{1}{2}\|x^t - x^*\|^2 - \frac{1}{2}\mathbb{E}_t \|x^{t+1} - x^t\|^2 + \eta(P_1 + P_2). \quad (25)$$

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

$$P_1 \leq \frac{\sigma_t^2 + D^2 \|x^t - x^*\|^2}{2\Delta_1} + \frac{\Delta_1 \mathbb{E}_t \|x^{t+1} - x^t\|^2}{2}. \quad (26)$$

Next, we decompose P_2 as

$$P_2 = \langle \mu^t - G(x^t), x^* - x^t \rangle + \mathbb{E}_t \langle \mu^t - G(x^t), x^t - x^{t+1} \rangle + \mathbb{E}_t \langle G(x^t) - G(x^{t+1}), x^* - x^{t+1} \rangle. \quad (27)$$

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

$$\langle \mu^t - G(x^t), x^* - x^t \rangle \leq B \|x^t - x^*\|^2. \quad (28)$$

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

$$\begin{aligned} \mathbb{E}_t \langle \mu^t - G(x^t), x^t - x^{t+1} \rangle &\leq \frac{\Delta_2 \|\mu^t - G(x^t)\|^2}{2} + \frac{\mathbb{E}_t \|x^t - x^{t+1}\|^2}{2\Delta_2} \\ &\leq \frac{\Delta_2 B^2 \|x^t - x^*\|^2}{2} + \frac{\mathbb{E}_t \|x^t - x^{t+1}\|^2}{2\Delta_2} \end{aligned} \quad (29)$$

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

$$\begin{aligned} \mathbb{E}_t \langle G(x^t) - G(x^{t+1}), x^* - x^{t+1} \rangle &\leq \frac{\Delta_3 \|G(x^t) - G(x^{t+1})\|^2}{2} + \frac{\mathbb{E}_t \|x^* - x^{t+1}\|^2}{2\Delta_3} \\ &\leq \frac{\Delta_3 L^2 \|x^t - x^{t+1}\|^2}{2} + \frac{\mathbb{E}_t \|x^* - x^{t+1}\|^2}{2\Delta_3} \end{aligned} \quad (30)$$

Putting together all the estimates (25)-(30) yields

$$\begin{aligned} \frac{1 + 2\eta\alpha - 2\eta\Delta_3^{-1}}{2} \mathbb{E}_t \|x^{t+1} - x^*\|^2 &\leq \frac{1 + \eta D^2 \Delta_1^{-1} + 2\eta B + \eta \Delta_2 B^2}{2} \|x^t - x^*\|^2 \\ &\quad - \frac{1 - \eta \Delta_1 - \eta \Delta_2^{-1} - \eta \Delta_3 L^2}{2} \mathbb{E}_t \|x^{t+1} - x^t\|^2 + \frac{\eta \sigma_t^2}{2\Delta_1}. \end{aligned} \quad (31)$$

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 (31) 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 (22).

We next prove the second claim. To this end, suppose $B = 0$. All of the reasoning leading up to and including (26) is valid. Continuing from this point, using Young's inequality, we upper bound P_2 by:

$$P_2 \leq \frac{\mathbb{E}_t \|\mu^t - G(x^{t+1})\|^2}{2\Delta_2} + \frac{\Delta_2 \mathbb{E}_t \|x^{t+1} - x^*\|^2}{2}. \quad (32)$$

Next observe

$$\begin{aligned} \mathbb{E}_t \|\mu^t - G(x^{t+1})\|^2 &\leq 2\mathbb{E}_t \|\mu^t - G(x^t)\|^2 + 2\mathbb{E}_t \|G(x^t) - G(x^{t+1})\|^2, \\ &\leq 2C_t^2 + 2L^2 \|x^t - x^{t+1}\|^2 \end{aligned} \quad (33)$$

and therefore

$$P_2 \leq \frac{2C_t^2 + 2L^2 \|x^t - x^{t+1}\|^2}{2\Delta_2} + \frac{\Delta_2 \mathbb{E}_t \|x^{t+1} - x^*\|^2}{2} \quad (34)$$

Putting the estimates (25), (26), and (34) together yields:

$$\begin{aligned} \frac{1 + \eta(2\alpha - \Delta_2)}{2} \mathbb{E}_t \|x^{t+1} - x^\star\|^2 &\leq \frac{1 + \eta D^2 \Delta_1^{-1}}{2} \|x^t - x^\star\|^2 \\ &\quad + \frac{\eta \sigma_t^2}{2\Delta_1} + \frac{2\eta C_t^2 \Delta_2^{-1}}{2} - \frac{1 - 2\eta L^2 \Delta_2^{-1} - \eta \Delta_1}{2} \mathbb{E}_t \|x^{t+1} - x^t\|^2 \end{aligned} \quad (35)$$

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 (35) directly yields (23). In the case $B \equiv C_t \equiv 0$, instead of (33) 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. $\square$

F Technical results on sequences

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

Lemma 16. *Consider a sequence $D_t \geq 0$ for $t \geq 1$ and constants $t_0 \geq 0$, $a > 0$ satisfying*

$$D_{t+1} \leq \left(1 - \frac{2}{t+t_0}\right) D_t + \frac{a}{(t+t_0)^2} \quad (36)$$

Then the estimate holds:

$$D_t \leq \frac{\max\{(1+t_0)D_1, a\}}{t+t_0} \quad \forall t \geq 1. \quad (37)$$

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

Lemma 17. *Consider a sequence $D_t \geq 0$ for $t \geq 1$ and constants $t_0 \geq 0$, $a, b > 0$ satisfying*

$$D_{t+1} \leq \left(1 - \frac{2}{t+t_0}\right) D_t + \frac{a}{(t+t_0)^2} + \frac{b}{(t+t_0)^3}. \quad (38)$$

Then the estimate holds:

$$D_t \leq \frac{\max\{(1+t_0)D_1/2, a\}}{t+t_0} + \frac{\max\{(1+t_0)^{3/2}D_1/2, b\}}{(t+t_0)^{3/2}} \quad \forall t \geq 1. \quad (39)$$

Proof. Clearly the recursion (38) 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 (39) holds automatically for the base case $t = 1$. As an inductive assumption, suppose the claim (39) 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. $\square$

G 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 18. 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 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

$$\begin{aligned} \frac{1}{2} \mathbb{E}[\|B^+ - V\|_F^2 \mid \mathcal{G}_2] &= \frac{1}{2} \|B - V\|_F^2 - \nu \langle B - V, (By - \mathbb{E}[b \mid \mathcal{G}_2])y^\top \rangle + \frac{\nu^2}{2} \mathbb{E}[\|(By - b)y^\top\|_F^2 \mid \mathcal{G}_2] \\ &= \frac{1}{2} \|B - V\|_F^2 - \nu \|(B - V)y\|_F^2 + \frac{\nu^2}{2} \|y\|^2 \mathbb{E}[\|By - b\|_F^2 \mid \mathcal{G}_2]. \end{aligned} \tag{40}$$

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 (40) 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

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

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. $\square$

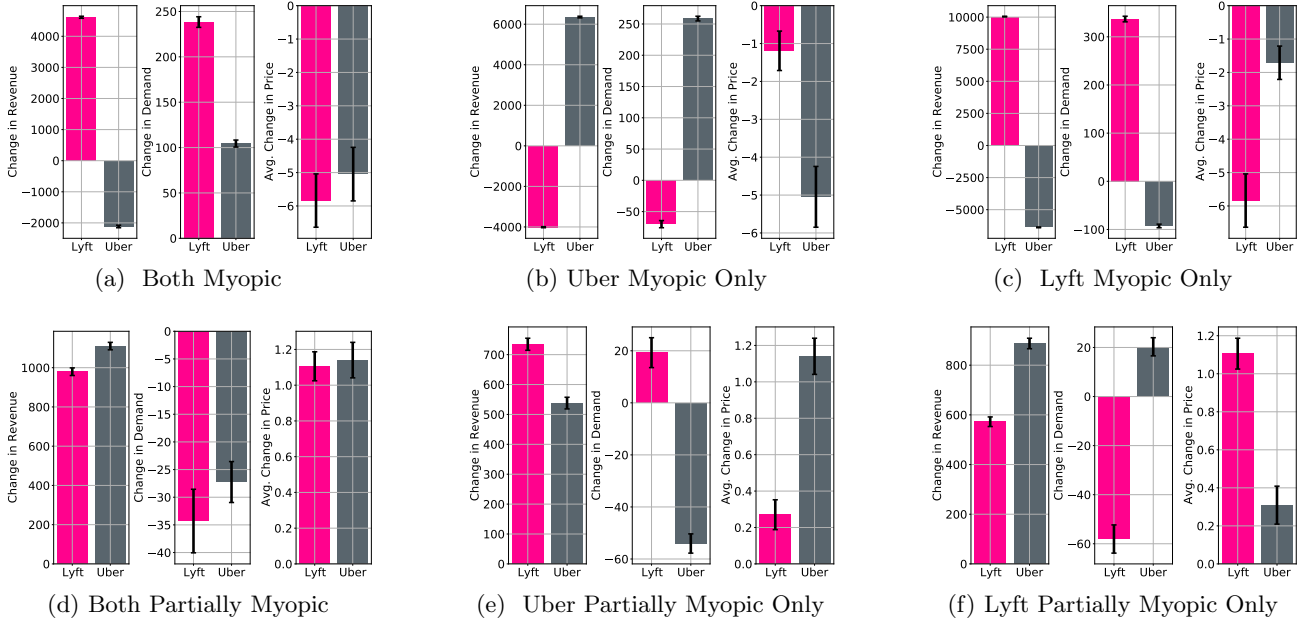


Figure 2: **Competition in Ride-Share Markets: Experiment 3.** Effects of players being (a)–(c) myopic or (d)–(f) partially myopic relative to Nash (not myopic, and consider competition). Positive changes in revenue indicate the Nash equilibrium is better for that player. When a player is myopic, they do not consider any performative effects in their update—i.e., $g_i^t = (\lambda_i I)^\top x_i^t - \frac{1}{2}\zeta_i^t$ —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.

H Additional Numerical Experiments

In this appendix, we describe the data construction for the semi-synthetic experiments and also provide additional experiments on the effects of competition.⁴

H.1 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. 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 $5 \cdot \lfloor \frac{p}{5} \rfloor$ where p is the price of the ride. Then, for each bin j we have a different base empirical distribution $\mathcal{P}_{i,j}$ for each player $i \in \{1, 2\}$ which is just the collection of rides for that bin.

For each bin, we estimate these price elasticity matrices A_i and A_{-i} from the data using the heuristic that a 50% increase in price by any firm leads to a 75% decrease in demand. With this heuristic we use the average base demand for each location and price bin pair to estimate both the diagonal elements of A_i and A_{-i} . In the experiments presented, our semi-synthetic model is such that there is no correlation between locations; 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. We further note that the results presented in this section are for the \$10 nominal price bin, however, in the repository of code it is easy to adjust this parameter to any of the other price bins. The conclusions we draw are similar across the bins.

H.2 Effects of Competition on Market Outcomes

Experiment 3: Effect of Ignoring Performativity. We study the impact of players ignoring performative effects due to competition in the data distribution. In Figure 2, we explore the effects of players either being completely myopic—i.e., $g_i^t = (\lambda_i I)^\top x_i^t - \frac{1}{2}\zeta_i^t$ —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. To compute the equilibrium outcomes we run the stochastic gradient method with a constant step-size of $\eta = 0.001$.

In Figure 2 (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. Interestingly, the player that is worse off at the Nash is the non-myopic player. In Figure 2 (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 2 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 2 (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.

⁴The code for the examples can be found <https://github.com/ratliffjlj/performativepredictiongames.git>.