
Emergent specialization from participation dynamics and multi-learner retraining

Sarah Dean
Cornell
University

Mihaela Curmei
University of California
Berkeley

Lillian J. Ratliff
University of
Washington

Jamie Morgenstern
University of
Washington

Maryam Fazel
University of
Washington

Abstract

Numerous online services are data-driven: the behavior of users affects the system’s parameters, and the system’s parameters affect the users’ experience of the service, which in turn affects the way users may interact with the system. For example, people may choose to use a service only for tasks that already works well, or they may choose to switch to a different service. These adaptations influence the ability of a system to learn about a population of users and tasks in order to improve its performance broadly. In this work, we analyze a class of such dynamics—where users allocate their participation amongst services to reduce the individual risk they experience, and services update their model parameters to reduce the service’s risk on their current user population. We refer to these dynamics as *risk-reducing*, which cover a broad class of common model updates including gradient descent and multiplicative weights. For this general class of dynamics, we show that asymptotically stable equilibria are always segmented, with sub-populations allocated to a single learner. Under mild assumptions, the utilitarian social optimum is a stable equilibrium. In contrast to previous work, which shows that repeated risk minimization can result in representation disparity and high overall loss with a single learner (Hashimoto et al., 2018; Miller et al., 2021), we find that repeated myopic updates with multiple learners lead to better outcomes. We illustrate the phenomena via a simulated example initialized from real data.

1 INTRODUCTION

Many online platforms, including social media networks, personalized recommendation engines, and advertising auction systems, collect user data and make incremental adjustments to the models they use to personalize content. These continuous updates are motivated by many factors, though large amongst them is the fact that the systems operate in non-stationary environments, where the preferences of their users change as the system operates. Changes in user preferences might occur exogeneously of service settings (e.g., global events might spur interest in new topics) or endogeneously (e.g., increasing the ranking of certain content on a platform might lead the content to “go viral”). The fact that user behavior might depend on service settings can take on many forms: people may learn to ignore or avoid clicking on advertisements; they may choose to use the service only for tasks at which it already works well; or they may choose to switch to a different service if they have a better experience with the second service. The latter two examples of adaptation, where users might opt for services that already suit their needs, affect the system’s capacity to learn about its user base and improve its overall performance.

In this work, we study a particular form of endogenously shifting distributions over multiple rounds, in contexts where individuals prefer to use services whose predictions are more accurate for them. Much of the existing work on endogenous distribution shift focuses on users who modify their features to achieve desired outcomes, as in strategic classification (Hardt et al., 2016) and related problems (Perdomo et al., 2020; Miller et al., 2021). While important, this model of data manipulation does not capture the most straightforward way that individuals express their preferences in a market: by choosing amongst alternative providers. In fact, recent work has shown that in the presence of a choice of participation between competing providers, individuals do not have an incentive to perform costly data manipulations (Hardt et al., 2022).

Consider as an example a social media platform. If the platform recommends content that does not appeal to the tastes of younger generations, these users will spend a smaller frac-

tion of their time on that platform. This results in positive (i.e., self-reinforcing) feedback loop, where a service’s poor performance on young customers dissuades them from using the service, leading to less data and diminishing weight placed on making better predictions for young customers in the future. Within a single service, these effects may lead to representation disparity (Hashimoto et al., 2018).

However, in a broader ecosystem, individuals can choose *amongst* services. If a new social media platform can predict the tastes of younger users more accurately, the younger users may spend more of their time on the new service, and correspondingly less on an existing platform. The new platform will then receive more data and improve its performance on young customers, while the old platform’s predictions may deteriorate, reinforcing their exit. Similarly, in the context of Large Language Models (LLM), if one LLM performs particularly well on creative tasks and another on answering homework questions, the distribution of prompts each receives may shift towards their existing expertise. Such feedback loops can also arise in settings such as music recommendation or healthcare, where demographic and socio-economic factors explain some of the emerging specialization (see examples in Appendix A).

In this paper, we study the dynamics of populations apportioning themselves amongst services, and services that choose predictors based on their observed user population. Our first contribution is to introduce and formalize this general setting. In Section 3 we introduce *risk reducing* populations and services who choose their actions myopically, incrementally improving their utility based on current conditions. Our second contribution is to present a complete characterization of stable fixed points for this general class of dynamics in Section 4. By drawing a connection between the dynamics and the *total risk*, our third contribution is to characterize the implications of this dynamic in terms of a utilitarian notion of *social welfare*, and argue that increasing the number of available services leads to better outcomes in terms of accurate predictions and user experience. In Section 5 we illustrate our theory with simulated experiments and conclude with a discussion of future work in Section 6.

2 RELATED WORK

The study of equilibria in the presence of utility optimizing agents has classical roots in game theory, and optimization over decision-dependent probabilities is classically studied by stochastic optimization and control (e.g., the review article by Hellemo et al. (2018)); we narrow our focus to the most relevant literature on this as it arises in machine learning systems.

Endogenous Distribution Shifts. In the study of machine learning systems, a large body of literature studies exogenous distribution shifts such as covariate, label, or concept drift (Quiñonero-Candela et al., 2008). A more recent trend

is to study shifts in the underlying data distribution due to endogenous reactions, for example due to strategic behavior exhibited by a user population. The work of Perdomo et al. (2020) introduces *performative prediction* as a model capturing user reaction via endogenous distribution shifts. This work models a single decision-maker facing a risk minimization problem subject to an underlying decision-dependent data distribution. Following its introduction, several relevant solution concepts have been explored and algorithms for achieving them proposed (Izzo et al., 2021; Drusvyatskiy and Xiao, 2020; Mandler-Dünner et al., 2020; Miller et al., 2021). A variant of the single decision-maker performative prediction problem studies time-dependent dynamics of the data distribution, with both exogenous (Wood et al., 2021; Cutler et al., 2021) and endogenous (Ray et al., 2022; Brown et al., 2022) sources. These works primarily consider strategic covariate shifts in a single distribution. In contrast, we consider a mixture of distributions: sub-populations of users whose participation choices result in attrition and retention dynamics which are not studied in the aforementioned distribution shift literature.

Multiple Decision-Makers. Endogenous distribution shift has also been studied in settings with multiple decision-makers as a continuous game. For instance, the multi-player performative prediction problem extends the original problem by allowing for multiple competing decision-makers (Narang et al., 2022; Piliouras and Yu, 2022; Wood and Dall’Anese, 2022). This line of work differs from ours in that the population is modeled as homogeneous and stateless. These works focus on characterizing the existence and uniqueness of different types of competitive equilibria for the game, and analyze learning dynamics that lead to different equilibrium concepts. In contrast, in our paper the focus is on asymptotically stable points (equilibrium) for the combined dynamical system resulting from myopic optimization by non-anticipating decision-makers and stateful user participation updates.

Retention. User retention in machine learning systems is closely related to the population participation dynamics we consider (Hashimoto et al., 2018; Zhang et al., 2019). In settings with multiple sub-populations of users of different types, the question of retention has been explored in parallel with the issue of fairness. Hashimoto et al. (2018) coined the term *representation disparity* for the phenomenon in which the traditional approach of minimizing average performance leads to high overall accuracy coupled with low accuracy on minority groups, causing an exodus of said groups. For single learners, systems which instead perform robust risk minimization avoid such disparity.

Our work generalizes the single-learner retention setting and analyzes the fixed points of dynamics between multiple systems and populations without modifying risk functions to be robust. Ginart et al. (2021) also consider user choice between multiple learning systems, with an empirical inves-

tigation and theoretical results in restricted settings focused on finite sample effects. In contrast, we propose a general class of *risk reducing* dynamics and develop a comprehensive theoretical understanding.

3 FRAMEWORK AND SETTING

We consider a setting where the population of individuals is composed of n subpopulations spreading their participation amongst m learners (service providers or decision-makers). Figure 1 illustrates a simple example. Each subpopulation $i \in \{1, \dots, n\} =: [n]$ has features and labels distributed according to a fixed distribution $(x, y) =: z \sim \mathcal{D}_i$ and makes up β_i proportion of the total population, so that $\sum_{i=1}^n \beta_i = 1$. An α_{ij} proportion of subpopulation i is associated to each learner $j \in [m]$, normalized so that $\sum_{j=1}^m \alpha_{ij} = 1$. The subpopulations therefore redistribute their participation among the various learners. Further, to model the ability of subpopulations to opt-out, one can include a static “null learner”.

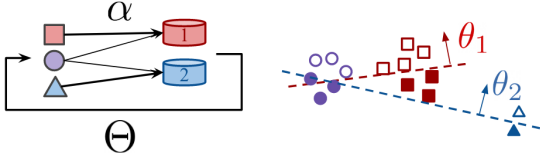


Figure 1: $n = 3$ subpopulations ($\square$, $\circ$, $\triangle$) select among $m = 2$ learners (red, blue) based on classification accuracy with respect to label (solid, hollow). Parameters $\Theta = (\theta_1, \theta_2)$ (decision lines) update in response to subpopulations participation $\alpha_{i,j}$. At the current state, the circle subpopulation will shift participation towards blue learner.

A subpopulation can be as broad as a demographic or affinity group and as granular as a single individual. The allocation of a subpopulation can represent several things: the fraction of a subpopulation which uses a given service, or the *fraction of time* users from that subpopulation choose to spend using learners’ systems. Accordingly, the relative size β_i of the population can represent the proportion of individuals or the *total time* individuals spend. This framework also allows for a subpopulation to represent *types of tasks or activities* a user wishes to accomplish, allocating these tasks to learners based upon which systems perform best on which tasks. *The only assumption we make about any subpopulation is that individual samples comprising it are i.i.d.*

Throughout, we assume that there are fewer learners than subpopulations, $m \leq n$. Each learner j observes data from the subpopulations who participate in it. Formally it observes features and labels drawn from the mixture distribution determined by the participation and subpopulation sizes:

$$(x, y)_j = z_j \sim \frac{\sum_{i=1}^n \alpha_{ij} \beta_i \mathcal{D}_i}{\sum_{i=1}^n \alpha_{ij} \beta_i}$$

Learners make predictions or decisions according to a parameter $\theta_j \in \mathbb{R}^d$. Beyond the information encoded in the features and labels, the learners are unaware of which subpopulation individual data points are.

The quality of predictions made by parameter $\theta_j \in \mathbb{R}^d$ for an individual instance z_j is quantified by the loss $\ell(\theta_j; z_j)$. The quality of θ_j for a subpopulation is quantified by the average loss, i.e. the *risk* $\mathcal{R}_i(\theta_j) = \mathbb{E}_{z \sim \mathcal{D}_i}[\ell(\theta_j; z)]$. Throughout, we will make the additional assumption that the risk function for each subpopulation $\mathcal{R}_i(\theta)$ is convex and differentiable. Figure 2 illustrates an example of the risk functions arising in linear regression.

3.1 Decision dynamics of learners and subpopulations

Subpopulations and learners react to each other; Updates in subpopulation allocations lead to updates in learners parameters $\Theta^t = (\theta_1^t, \dots, \theta_m^t)$, and vice versa. We introduce a broad class of update dynamics by way of a canonical example. Suppose that each subpopulation i updates its allocation by increasing the participation proportional to the quality of various models; for example, by spending more time on recommendation platforms that suggest more engaging content. Recalling that the risk (i.e. average loss) quantifies quality, this manifests as a *multiplicative weights update*: $\alpha_{ij}^{t+1} \propto \alpha_{ij}^t \cdot \exp(-\gamma \mathcal{R}_i(\theta_j))$ for $j \in [m]$ and some parameter $\gamma > 0$. This is similar to the retention function studied by Hashimoto et al. (2018) and has connections to replicator dynamics, a foundational evolutionary dynamic that can be interpreted as a process of information diffusion and imitation (Sandholm 2020).

Recall that each learner j observes data from the mixture distribution $(\sum_{i=1}^n \alpha_{ij} \beta_i)^{-1} \sum_{i=1}^n \alpha_{ij} \beta_i \mathcal{D}_i$ for which we use the shorthand $\mathcal{D}(\alpha_{:,j})$, where $\alpha_{:,j} \in \mathbb{R}^n$ denotes the vector of allocations from all subpopulations to learner j . Suppose the learners update their parameters using gradient descent to reduce the average loss over this data (e.g. to improve the prediction of user engagement). Setting aside finite sample issues, for a step size γ_t the gradient update takes the form $\theta_j^{t+1} = \theta_j^t - \gamma_t \nabla_{\theta} \mathbb{E}_{z \sim \mathcal{D}(\alpha_{:,j})}[\ell(\theta_j^t; z)]$. This is an incremental version of the *repeated retraining dynamics* which have been studied in the single learner setting by Hashimoto et al. (2018); Perdomo et al. (2020).

Despite the apparent simplicity of independent update rules, the evolution of subpopulations and learners is highly coupled. The sequential interaction between subpopulations and learners leads to complex nonlinear dynamics: i.e. multiplicative weights over non-stationary risks (due to learner updates) and gradient descent over non-stationary data distributions (due to subpopulation updates). To study this complex behavior, we now formalize key properties.

The first observation is that updates are *stateful*, with subpopulation allocations and learner parameters depend-

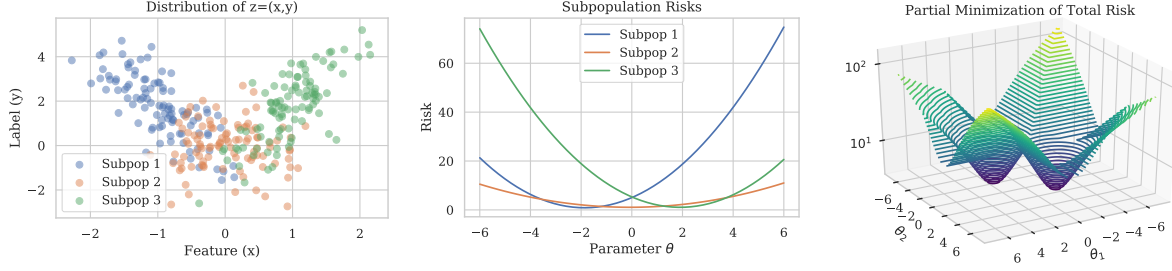


Figure 2: An example arising from least-squares linear regression with $n = 3$ subpopulations and $m = 2$ learners. Left: The distribution of $z = (x, y)$, colored by subpopulation. Middle: The subpopulation risks $\mathcal{R}_i(\theta)$ arising from least-squares linear regression $\ell(\theta; z) = (y - \theta x)^2$. Right: A visualization of the non-convex total risk as a function of learner parameters, via the partial minimization over subpopulation allocation: $\min_{\alpha} \sum_{i=1}^3 \sum_{j=1}^2 \alpha_{ij} \mathcal{R}_i(\theta_j) = \sum_{i=1}^3 \min\{\mathcal{R}_i(\theta_1), \mathcal{R}_i(\theta_2)\}$.

ing on previous values. This motivates a description of the dynamics arising from interactions between n subpopulations and m learners in terms of the overall state $\alpha \in \Delta_m \times \cdots \times \Delta_m =: \Delta_m^n$ and $\Theta \in \mathbb{R}^{m \times d}$. We thus define for each subpopulation i a general Markovian allocation function $\nu_i^t : \Delta_m \times \mathbb{R}^{m \times d} \rightarrow \Delta_m$ which describes the participation update $\alpha_{i,:}^{t+1} = \nu_i^t(\alpha_{i,:}^t, \Theta^t)$ at time t , where $\alpha_{i,:} \in \Delta_m$ denotes the vector of allocations from the subpopulation i to all learners. Similarly, define $\mu_j^t : \mathbb{R}^d \times \mathbb{R}^n \rightarrow \mathbb{R}^d$ so learner j updates their parameter according to $\theta_j^{t+1} = \mu_j^t(\theta_j^t, \alpha_{:,j}^t)$.

The second observation is that the basis for the updates is the average loss, i.e. risk. This motivates the following definition: given participation α and parameters Θ , the average risk experienced by each subpopulation i and each learner j is:

$$\begin{aligned} \bar{\mathcal{R}}_i^{\text{subpop}}(\alpha_{i,:}, \Theta) &:= \mathbb{E}_{j \sim \alpha_{i,:}} \left[\mathbb{E}_{z \sim \mathcal{D}_i} [\ell(\theta_j; z)] \right], \\ \bar{\mathcal{R}}_j^{\text{learner}}(\alpha_{:,j}, \theta_j) &:= \mathbb{E}_{z \sim \mathcal{D}(\alpha_{:,j})} [\ell(\theta_j; z)]. \end{aligned}$$

In the recommendation example, $\bar{\mathcal{R}}^{\text{subpop}}$ captures the dissatisfaction with content for a subpopulation and $\bar{\mathcal{R}}^{\text{learner}}$ corresponds to the average prediction error of the platform. Intuitively, multiplicative weights reduces the average subpopulation risk while gradient descent reduces the average learner risk.

Definition 3.1 (Reducing and Minimizing Dynamics). A u update rule is P -reducing w.r.t. v if $P(u^{t+1}, v^t) \leq P(u^t, v^t)$ for all t and any sequence of v^t . It is further P -minimizing in the limit if the inequality is strict when u^t is not a minimizer and $\lim_{t \rightarrow \infty} P(u^t, v) = \min_u P(u, v)$.

We call a subpopulation i *risk reducing* (resp. minimizing) when the allocation update on $\alpha_{i,:}$ is $\bar{\mathcal{R}}_i^{\text{subpop}}$ -reducing (resp. minimizing in the limit) with respect to Θ . Similarly, we call a learner j *risk reducing* (resp. minimizing) when the parameter update on θ_j is $\bar{\mathcal{R}}_j^{\text{learner}}$ -reducing (resp. minimizing in the limit) with respect to $\alpha_{:,j}$.

We remark that the notion of risk minimizing in the limit is reasonable for subpopulations because their average risk is linear in $\alpha_{i,:}$. It is also reasonable for learners because their average risk is convex in θ_j (due to the assumption that risks $\mathcal{R}_i(\theta_j)$ are convex). However, risk-reducing/minimizing is only a property defined with respect to the participation α or parameter Θ observed at a previous time step. Thus it does not necessarily hold that $\bar{\mathcal{R}}_j^{\text{learner}}$ or $\bar{\mathcal{R}}_i^{\text{subpop}}$ decrease when the state evolves $(\alpha^t, \Theta^t) \rightarrow (\alpha^{t+1}, \Theta^{t+1})$ by sequential updates of ν^t and μ^t . Our experiments (Figure 4a) illustrate the non-monotonicity of the coupled updates.

Example 3.2 (Semi-static participation). Suppose a population has a constant allocation of 20% to one learner, while the remaining 80% is allocated to the remaining learners inversely proportional to the learner's risk on that population. This is risk reducing but not risk minimizing in the limit.

Example 3.3 (Full risk minimization). Suppose that a learner updates its parameter to minimize the average risk function $\bar{\mathcal{R}}_j^{\text{learner}}(\alpha_{:,j}^t, \cdot)$ at each timestep. This has been studied as *repeated retraining dynamics* in the single learner case by Hashimoto et al. (2018); Perdomo et al. (2020).

Proposition 3.4. A subpopulation i updating their participation with multiplicative weights is risk minimizing in the limit if $\gamma > 0$ and $\alpha_{ij}^0 > 0 \forall j$. A learner updating its parameter with gradient descent is risk minimizing in the limit when the risk functions $\mathcal{R}_i(\theta)$ are L smooth and the step size satisfies $\gamma^t < \frac{2}{L}$, $\sum_{t=0}^{\infty} \gamma^t = \infty$, and $\sum_{t=1}^{\infty} (\gamma^t)^2 < \infty$.

We provide a proof and detail several other examples of risk reducing dynamics in Appendix D.1

3.2 Equilibria and stability

We focus on the equilibrium states resulting from risk-reducing subpopulations and learners.

Definition 3.5 (Equilibrium). The state $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is an equilibrium state if it is stationary under the dynamics update $\{\nu_i^t\}, \{\mu_j^t\}$; i.e. that for all $i \in [n]$ and $j \in [m]$:

$$\alpha_{i,:}^{\text{eq}} = \nu_i^t(\alpha_{i,:}^{\text{eq}}, \Theta^{\text{eq}}) \quad \text{and} \quad \theta_j^{\text{eq}} = \mu_j^t(\alpha_{:,j}^{\text{eq}}, \theta_j^{\text{eq}}).$$

If learners and subpopulations are in an equilibrium state, they will remain that way indefinitely. However, some equilibrium configurations may be unstable to perturbations.

Definition 3.6 (Stable Equilibrium). The state $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is a *stable* equilibrium state if it is an equilibrium and for each $\epsilon_\alpha, \epsilon_\theta > 0$, there exist $\delta_\alpha, \delta_\theta > 0$ such that

$$\begin{aligned} \|\alpha^0 - \alpha^{\text{eq}}\| < \delta_\alpha, \implies \|\alpha^t - \alpha^{\text{eq}}\| &\leq \epsilon_\alpha, \forall t \geq 0. \\ \|\Theta^0 - \Theta^{\text{eq}}\| < \delta_\theta \implies \|\Theta^t - \Theta^{\text{eq}}\| &\leq \epsilon_\theta \end{aligned}$$

It is further *asymptotically stable* if $\lim_{t \rightarrow \infty} \|\alpha^t - \alpha^{\text{eq}}\| = 0$ and $\lim_{t \rightarrow \infty} \|\Theta^t - \Theta^{\text{eq}}\| = 0$.

Stability analysis identifies qualitatively different equilibrium states. For the class of risk reducing dynamics that we study, equilibria may be unstable, stable, or asymptotically stable; Appendix D.2 presents examples. While a quantitative understanding of convergence may also be of interest, it would require stronger assumptions on the behavior of subpopulations and learners; here we favor generality and leave this to future work. Furthermore, characterizing stable equilibria sets the foundation for understanding high probability behavior of systems under noisy updates which are risk reducing only in expectation (Kushner, 1967). This sets the stage for finite sample risk minimization or multi-agent user models, a challenge which we leave to future work.

4 MAIN RESULTS

We study a large class of feedback dynamics between risk reducing learners and subpopulations described by the sequential updates: $\alpha^{t+1} = \nu^t(\alpha^t, \Theta^t)$ and $\Theta^{t+1} = \mu^t(\alpha^{t+1}, \Theta^t)$. Our analysis allows for learners and subpopulations who exhibit a diverse range of behaviors. We do not require that every learner or every subpopulation update their parameter or allocation in the same manner or even at every timestep, allowing for any number of round-robin schemes. Our only assumption on learner and subpopulation updates is that they are risk reducing or minimizing.

Figure 3 presents a summary of the equilibria characterization that we present in this section. All omitted proofs can be found in Appendix C.

4.1 Total Risk Reduction

Definition 4.1 (Total Risk). The *total risk* of all subpopulations over all learners is the weighted sum

$$\mathcal{R}^{\text{total}}(\alpha, \Theta) := \sum_{i=1}^n \sum_{j=1}^m \beta_i \alpha_{ij} \mathcal{R}_i(\theta_j).$$

The total risk maps $\Delta_m^n \times \mathbb{R}^{m \times d} \rightarrow \mathbb{R}$. While our assumption that the loss is convex implies that the total risk is convex in Θ , it is not jointly convex in (α, Θ) , illustrated in the right panel of Figure 2.

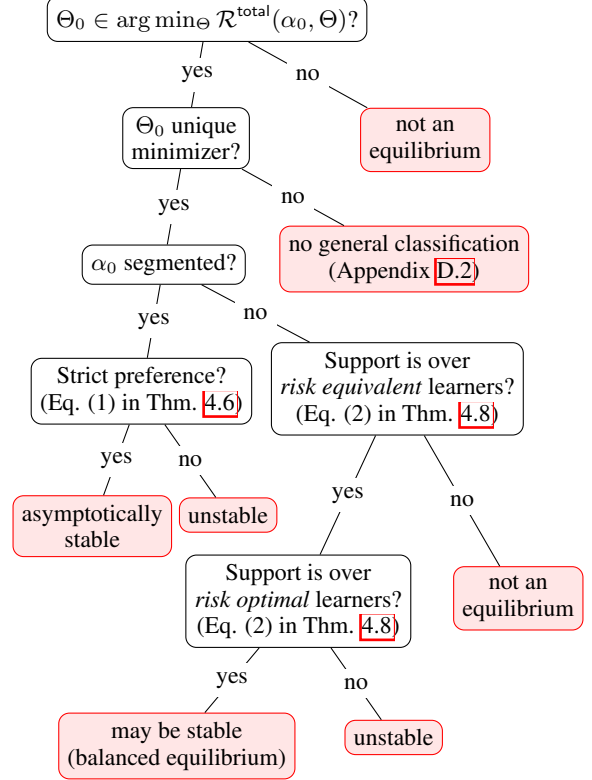


Figure 3: A summary of our main results on equilibria classification for a given participation α_0 and model parameters Θ_0 . These results hold for dynamics which are risk minimizing in the limit and loss functions that are convex.

Our first result shows that the total risk $\mathcal{R}^{\text{total}}(\alpha^t, \Theta^t)$ is non-increasing over time.

Proposition 4.2. *For any risk-reducing subpopulation and learner dynamics, the total risk is non-increasing: $\mathcal{R}^{\text{total}}(\alpha^{t+1}, \Theta^{t+1}) \leq \mathcal{R}^{\text{total}}(\alpha^t, \Theta^t)$, $\forall t$. If subpopulations and learners are risk minimizing in the limit, then the total risk is strictly decreasing unless (α^t, Θ^t) is a local minimizer of $\mathcal{R}^{\text{total}}$.*

Proof Sketch. First note that the total risk can be decomposed into either a weighted sum of average subpopulation risk or average learner risk. Thus the fact that learner and subpopulation dynamics are risk reducing ensures that the total risk is decreasing after the sequential updates. $\square$

Thus, the total risk acts like a potential function for the feedback dynamics of learners and subpopulations. When the subpopulation and learner dynamics are risk minimizing in the limit, there is a strong connection between properties of the total risk function and equilibria of the dynamics.

Theorem 4.3. *For any learners and subpopulations who are risk minimizing in the limit, an equilibrium $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is asymptotically stable if it is an isolated local minimizer*

of the total risk $\mathcal{R}^{\text{total}}$. If it is not a local minimizer of the total risk, then it is not stable.

Proof Sketch. By Proposition 4.2 the function $V(\alpha, \Theta) := \mathcal{R}^{\text{total}}(\alpha, \Theta) - \mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is potential function for the autonomous dynamical system $(\alpha^t, \Theta^t) \rightarrow (\alpha^{t+1}, \Theta^{t+1})$. The stability result follow from Lyapunov arguments. $\square$

The connection between stability and the total risk function is significant in at least two ways: first, it means that under general classes of myopic and self-interested behaviors on the part of subpopulations and learners, the total risk is driven to at least a local minimum. Second, it is a technically useful connection that will enable us to characterize and classify the stable equilibria for dynamics which are risk minimizing in the limit. We remark that Theorem 4.3 leaves open the question of stability for equilibria which are non-isolated minima of the total risk function. In Appendix D.2, we provide examples which show that such points may be asymptotically stable, stable, or unstable depending on the particular instantiation of dynamics. The following existence result further motivates our focus dynamics which are risk minimizing, rather than just reducing.

Corollary 4.4. *Equilibria exist when learners and subpopulations are risk minimizing in the limit and the total risk function has isolated local minima. They may not exist otherwise.*

Example of dynamics without equilibria. Consider subpopulations with risk functions minimized at the same value θ^* . If learners use full risk minimization, the setting lacks isolated minima because the total risk is uniform across all allocations α . Assuming that risk-minimizing subpopulations randomly choose among equivalent learners, no equilibrium exists as allocations randomly switch between learners once the learners converge to the optimum θ^* . $\square$

4.2 Segmented and Balanced Equilibria

Definition 4.5 (Segmented allocation). An allocation is segmented if $\alpha_{ij} \in \{0, 1\}$ for all i, j .

In a segmented allocation, each subpopulation is associated with a single learner, and thus the population is partitioned across learners. For allocation dynamics like multiplicative weights, such configurations are clearly equilibria for any parameter choice Θ on the part of the learners. We thus consider the set of possible segmented equilibria and characterize which are asymptotically stable.

Theorem 4.6. *Suppose learners and subpopulations are risk minimizing in the limit, α^{eq} is segmented, and $\mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta)$ has a unique minimizer Θ^{eq} . Define a mapping $\gamma : [n] \rightarrow [m]$ such that $\gamma(i) = j$ is the learner with nonzero mass in $\alpha_{i,\cdot}^{\text{eq}}$. If every subpopulation strictly prefers*

their current learner:

$$\mathcal{R}_i(\theta_{\gamma(i)}^{\text{eq}}) < \mathcal{R}_i(\theta_j^{\text{eq}}), \quad (1)$$

for all i and learners $j \neq \gamma(i)$, then $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is an asymptotically stable equilibrium. If there is a subpopulation who would strictly prefer to switch learners, then $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is not stable.

When risks are strongly convex, there is always such a unique minimizer Θ^{eq} . In particular, in a segmented allocation, each θ_j^{eq} minimizes the average loss over the group of subpopulations assigned to them.

Corollary 4.7. *Suppose that risk functions satisfy $\mathcal{R}_i(\theta) < \mathcal{R}_i(\theta') \iff \|\theta - \phi_i\| < \|\theta' - \phi_i\|$ for ϕ_i the subpopulation optimal parameter. Then in an asymptotically stable segmented equilibrium, the convex hulls of the grouped subpopulations optimal parameters $\{\phi_i\}$ are non-intersecting.*

Proof Sketch. Consider a partition where the convex hulls intersect for some pair of learners. Then there exists at least one subpopulation who would be better off switching to the other learner, and thus the risk condition in Theorem 4.6 cannot hold. $\square$

Applying the Corollary to the example in Figure 2, we see that a segmented equilibrium with subpopulation 1 and 3 participating in the same learner cannot be stable.

Theorem 4.6 leaves open the question of stability in the case that the risks in Equation (1) are equal. Under such *risk equivalence*, is it natural to consider equilibria where a subpopulation has support over multiple learners.

Theorem 4.8. *Consider dynamics which are risk minimizing in the limit and an α^{eq} with any subpopulation i having nonzero support on set of two or more learners $j \in \mathcal{J}$. Assume risks are strongly convex and define $\Theta^{\text{eq}} = \arg \min \mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta)$. Then $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ cannot be stable unless it is “balanced” in the sense that learners in $\mathcal{J}$ are risk equivalent and optimal for i , i.e. for all $j, j' \in \mathcal{J}$,*

$$\mathcal{R}_i(\theta_j^{\text{eq}}) = \mathcal{R}_i(\theta_{j'}^{\text{eq}}) \quad \text{and} \quad \nabla \mathcal{R}_i(\theta_j^{\text{eq}}) = 0. \quad (2)$$

If it is balanced, so are all allocations for subpopulation i with support over $\mathcal{J}$. Finally, all stable equilibria must be either balanced or segmented.

This result characterizes a set of possibly stable equilibria. It demonstrates that risk *optimality*, in addition to *equivalence*, is necessary. Guaranteeing the stability of such balanced equilibria requires further information about the dynamics, and it is not possible to make a general statement. Examples in Appendix D.2 demonstrate that such balanced equilibria may be asymptotically stable, stable, or unstable. Furthermore, the balance condition is fragile in the sense that it would not hold under small perturbations to the underlying risk functions. While the number of possible balanced

equilibria is combinatorial in the number of learners and subpopulations, risk functions are continuous, so it is possible to find arbitrarily small perturbations to any the risk functions that would destabilize all balanced equilibria.

Proof Sketch of Theorems 4.6 and 4.8 By Theorem 4.3 characterizing the stable equilibria is equivalent to characterizing isolated and non-isolated local minima of the total risk. We show that it suffices to characterize local minima of the partial minimization $F(\alpha) = \min_{\Theta} \mathcal{R}^{\text{total}}(\alpha, \Theta)$ over the simplex product $\Delta_m^n = \Delta_m \times \cdots \times \Delta_m$. Since F is concave, all minima occur on the boundary, i.e. a face or a vertex. Since F is still concave when restricted to a face of the simplex, the same argument shows the minima are on the boundary, hence vertices, except for the degenerate case where F takes a constant value over the face.

Thus, the isolated local minima occur at vertices of the simplex product, which correspond to segmented allocation. Further analysis of F yields the conditions presented in Theorem 4.6. The local minima in the degenerate case are characterized by the balanced equilibria conditions in Theorem 4.8. $\square$

4.3 Social Welfare for Segmented Populations

Definition 4.9. The *social welfare* of a state (α, Θ) is strictly decreasing in the total risk $\mathcal{R}^{\text{total}}(\alpha, \Theta)$.

This definition of social welfare is utilitarian in the sense that it depends on the cumulative quality of individuals' experiences. Maximizing the social welfare corresponds to minimizing the total risk, which can be posed as the following optimization problem

$$\begin{aligned} (\alpha^*, \Theta^*) \in \arg \min_{\alpha, \Theta} \quad & \mathcal{R}^{\text{total}}(\alpha, \Theta) \\ \text{s.t.} \quad & \alpha_{i,:} \in \Delta_m \quad \forall i = \{1, \dots, n\}. \end{aligned} \quad (3)$$

Here, (α^*, Θ^*) is the social welfare maximizer.

Our discussion of stable equilibria has so far focused on only local minimizers of the total risk. In fact, global minimization of this objective (and therefore maximization of social welfare) is a hard problem. The total risk objective can be viewed as an instance of the k -means clustering problem with $k = m$. In the language of this literature (e.g., Selim and Ismail (1984)), each subpopulation is a data point and the parameter selected by each learner is a cluster center. The allocations described by α correspond to (fuzzy) cluster assignment and each risk function $\mathcal{R}_i(\theta_j)$ corresponds to a measure of “dissimilarity” between data points (subpopulations) and cluster centers (learners).

The connection to k -means clustering elucidates the difficulty of minimizing the total risk. The “minimum sum-of-squares clustering” problem (i.e., squared Euclidean norm

dissimilarity) is NP hard with general dimension even when $k = 2$ (Aloise et al., 2009). When the number of clusters and dimension are fixed, Inaba et al. (1994) present an algorithm for solving the minimum sum-of-squares clustering problem which is polynomial in the number of datapoints. Translated to our setting, its complexity is $O(n^{md})$. It is therefore unrealistic to hope that a myopic dynamic might generally lead to social welfare maximization. However, due to the connections with total risk, risk reducing dynamics are at least well-behaved with regards to social welfare.

Proposition 4.10. *For risk reducing subpopulations and learners, social welfare is non-decreasing over time. If the dynamics are furthermore risk minimizing in the limit, social welfare is strictly increasing and stable equilibria correspond to local social welfare maxima.*

Local maximization is not a panacea: Example 4.11 shows a local maximum of the social welfare can be much worse than the global one.

Example 4.11 (Arbitrarily high total risk at local optimum). Consider three subpopulations with

$$\mathcal{R}_1(\theta) = \theta^2, \quad \mathcal{R}_2(\theta) = (\theta - 1)^2, \quad \mathcal{R}_3(\theta) = (\theta - \phi)^2$$

for some $\phi > 2$. Suppose that subpopulation sizes are $\beta_1 = \beta_2 = \beta$ and $\beta_3 = 1 - 2\beta$ for some $0 < \beta < 1/2$. Further suppose that there are two learners. Up to permutation, the social welfare optimum is $\theta_1 = 1/2$ and $\theta_2 = \phi$, with total risk $\beta/2$. However, as long as $\phi < \frac{1-\beta}{1-2\beta}$, there is another stable equilibrium. Let $\phi = \frac{1-\beta}{1-2\beta} - \epsilon$. Then the following is a stable equilibrium: $\theta_1 = 0$ and $\theta_2 = 1 - \epsilon$. The total risk is $\beta + \frac{(\beta-\epsilon)^2}{1-2\beta}$. For β close to $1/2$, this risk can be arbitrarily larger than the social optimum.

In this example, a large gap between a stable local optimum and the global optimum arises in part due to a large difference in subpopulations' sizes. We further remark that minority groups can be under-served particularly when considering worst-case risk over subpopulations (Hashimoto et al., 2018). Even at a social welfare maximizer (α^*, Θ^*) , the worst-case subpopulation risk can be arbitrarily bad. It is straightforward to construct such examples even in the single learner case: consider a minority group with vanishingly small population proportion and arbitrarily high risk at the optimal parameter for the majority group (Example D.10).

Despite these inherent difficulties, we find that the situation improves as the number of learners increases. It is straightforward to see that the maximal social welfare will increase: any point which is optimal for m learners can be trivially transformed into a feasible point for $m + 1$ learners which achieves the same social welfare, by allocating no subpopulations to the new learner. There is more nuance involved when considering any possible stable equilibria. Instead, we make a statement about a particular learner growth process which corresponds to existing learner m “splitting in half”.

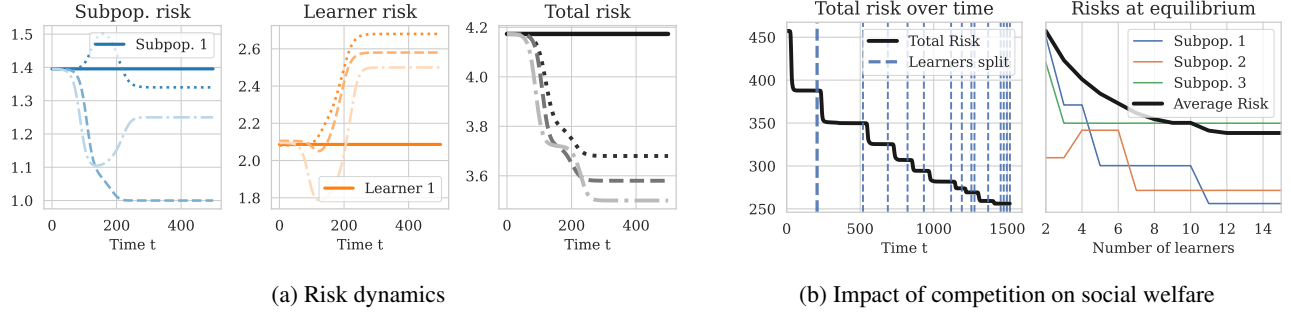


Figure 4: **Synthetic settings:** Figure (a) illustrates a setting with 3 subpopulations and 2 learners. The dsolid lines correspond to the risk trajectory for the unstable balanced equilibrium at initialization. Dotted and dashed lines illustrate risk trajectories under three different slight perturbations from the initialization. In Figure (b), the left plot illustrates the reduction in total risk over time. The dashed blue lines indicate when a new learner joins. The right plot shows the equilibrium-risk for a subset of the subpopulations as the number of learners increases.

Proposition 4.12. *Suppose that risks are strongly convex, there are m learners, $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is an equilibrium, and at least one subpopulation i allocated to learner m does not have optimal subpopulation risk, so $\nabla \mathcal{R}_i(\theta_m^{\text{eq}}) \neq 0$. The state is amended to add an additional learner: $\tilde{\Theta}^{\text{eq}} = [\Theta^{\text{eq}}, \theta_m^{\text{eq}}]$ and*

$$\tilde{\alpha}^{\text{eq}} = \begin{cases} \alpha_{:,j}^{\text{eq}} & j \leq m \\ \frac{1}{2} \alpha_{:,m}^{\text{eq}} & j \in \{m, m+1\} \end{cases}$$

Under dynamics which are risk minimizing in the limit, the equilibrium $(\tilde{\alpha}^{\text{eq}}, \tilde{\Theta}^{\text{eq}})$ is not stable, so a small perturbation will send the system to a state with strictly lower total risk (higher social welfare).

5 SIMULATIONS

We illustrate the salient properties of the decision dynamics in simulation¹. We consider both a synthetic setting as well as one instantiated from a prediction task on census data.

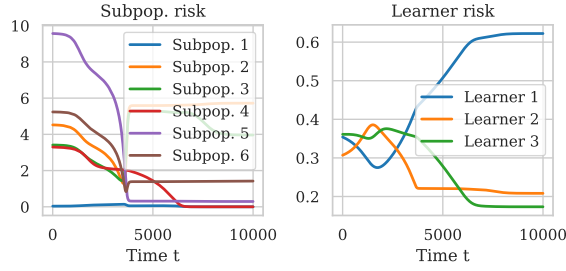
Synthetic In Figure 4a we consider a simple scenario with $n = 3$ subpopulations of equal sizes $\beta_i = 1/3$, quadratic risk functions $\mathcal{R}_i = \|\phi_i - \theta\|^2 + 1$ with distinct risk minimizing decisions ϕ_i and $m = 2$ learners. The learners minimize their risk according to *full risk minimization* (Example 3.3) and the subpopulations update their participation via multiplicative weights update (Section 3.1). When $\alpha_{i,j}^0 = 1/2$ for all i, j the risk equality condition from Theorem 4.8 is satisfied with $\theta_j^{\text{eq}} = (\phi_1 + \phi_2 + \phi_3)/3$, however the optimality condition is not. We therefore observe that this equilibrium is not stable, and slightly perturbing the initial conditions leads to split-market equilibria. Figure 4a illustrates trajectories from three different perturbations. It demonstrates that the total risk is non-increasing whereas the average risks for both learners and subpopulations are not monotonic. Each

of the perturbations has different risk trajectories and equilibrates at a different split-market equilibrium. We repeat these experiments with noisy dynamics, we consider both exogenous noise that independently perturbs the decisions of the learners and/or populations as well as intrinsic noise due to making updates with finite sample estimates rather than at population level. We find that the key properties of the dynamics hold when the updates are noisy, detailed experiments are presented in Appendix E.

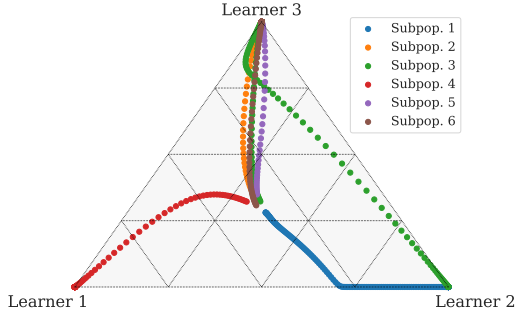
Another set of experiments in Figure 4b illustrates how a larger number of learners lead to better outcomes in terms of total risk. We consider a set of $m = 2$ learners and $n = 50$ subpopulations. We simulate the dynamics until the market has reached equilibrium, at which point a randomly chosen learner breaks up into two identical learners with half the user base. From this unstable equilibrium (Proposition 4.12) we slightly perturb the parameters of the two learners and allow the system to reach a new equilibrium state. The procedure repeats until the number of learners reaches number of subpopulations. These simulations illustrate that more competition improves social welfare, however the improvements are not uniform for all subpopulations with some groups seeing their risk at equilibrium increase with the addition of new learners.

Census data We consider a semi-synthetic setting where subpopulations and their risk functions are instantiated by a prediction task on real data. Using *folktables* (Ding et al., 2021) we consider a modified version of *ACSTravelTime* prediction problem derived from the 2018 California Census data. We consider 6 subpopulations corresponding to racial groups with relative size ranging from 1.2% to 61%. We define the least-squares risk functions as $\mathcal{R}_i(\theta) = \frac{1}{N_i} \|X_i \theta - y_i\|^2$ where $X_i \in \mathbb{R}^{N_i \times d}$ are the features (containing demographics, educational attainment, income levels, and modes of transportation) and $y_i \in \mathbb{R}^{N_i}$ are the labels (log transform of the daily commute time in minutes) for individuals within subpopulation i . We sim-

¹Implementation details and reproduction instructions at: <https://github.com/mcuermei627/MultiLearnerRiskReduction>



(a) Risk dynamics



(b) Allocation trajectories

Figure 5: Empirical subpopulations from Census data:

Figure (a) displays the relative risk with respect to the best achievable risk for the subpopulation over time. Figure (b) illustrates how allocations initialized near $(1/3, 1/3, 1/3)$ converge to a split market equilibrium.

ulate risk reducing dynamics from a perturbed balanced equilibrium over 3 learners. As in the synthetic example, the risks of learners and subpopulations are not all monotone (Figure 5a), but the total risk function is. Finally Figure 5b illustrates the convergence of allocation dynamics to a segmented equilibrium.

6 DISCUSSION

In this paper, we study the feedback dynamics of user retention for loss minimizing learners, where subpopulations choose between providers. We introduce a formal notion of *risk reducing* and *minimizing* to capture this feedback, and show that there is a close connection between such dynamics and the *total risk* summed over subpopulations and learners. We provide a comprehensive characterization of stable equilibria and investigate the implications in terms of a utilitarian social welfare. This work relates to questions of fairness and minority representation in several ways. First, our results imply that risk-minimizing dynamics in multi-learner settings can result in higher welfare for small subpopulations compared with single-learner settings, as studied by (Hashimoto et al., 2018; Zhang et al., 2019).

This resonates with recent work showing that monopolies have higher *performative power* and lead to lower individual utility (Hardt et al., 2022).

The dynamics that we study often lead to *segmentation* of subpopulations across learners as an emergent phenomenon². This segmentation can lead to pointwise lower risks for subpopulations, especially when subpopulations have considerably different risk profiles. In some contexts, the benefits of the reduced risk among subpopulations may outweigh possible harms from segregation. In others, where proportional representation of groups across learners, models, or clusters (Kleindessner et al., 2019a,b) is important, our work implies that independent risk minimization can lead to undesirable outcomes. In short, this work analyzes natural dynamics with consequences for the distribution of subpopulations amongst independent learners; whether or not the consequences are desirable depend on the specific application considered.

We highlight several directions for future work. Our results lay the groundwork for an investigation of the stochastic dynamics that occur for finite sample approximations to the risk or participation driven by decisions of individuals. Such behaviors are risk reducing in *expectation*, so we expect the noisy trajectories to converge with high probability to sets around the asymptotically stable equilibria we characterize. There are many interesting and relevant questions in the finite sample setting: What is the effect of sample size on the ability of new learners to enter a market and minority subpopulations to be adequately represented? Can we model heterogeneous learners who differ in which features they measure and with how much noise? Are there trade-offs between the expressivity of models and the practical difficulty of minimizing risk from finite samples in high dimensions?

It would also be interesting to consider extensions or alternative dynamics models for the learner and subpopulation decisions. One could investigate competitive learners who explicitly strategize to capture subpopulations (Ben-Porat and Tennenholtz, 2019; Aridor et al., 2020); this setting is related to facility location and Hotelling games (Owen and Daskin, 1998; Hotelling, 1929). One might imagine that subpopulations do not act uniformly and may not even be entirely independent of each other—the participation update may depend on some underlying social network. The connections between total risk reduction and k -means clustering algorithms suggest interventions such as subpopulation-aware initialization (Bose et al., 2023) that could improve social welfare. Results on “ground truth recovery” may yield insight into particular population structures that lead to simpler dynamics or restricted sets of equilibria.

²This connects to economic literature on “rational” discrimination, where competitors have no inherent preference to discriminate and yet equilibria are segregated, e.g. Foster and Vohra (1992)

Acknowledgements

We thank Laurent Lessard for suggesting the clever argument to prove Lemma B.7. JM is supported by an NSF CAREER award (ID 2045402), an NSF AI Center Award (The Institute for Foundations of Machine Learning), and the Simons Collaboration on the Theory of Algorithmic Fairness. LJR is supported by NSF CNS-1844729, NSF IIS-1907907, and Office of Naval Research YIP Award N000142012571. MF is supported by NSF TRIPODS II-DMS 2023166 and NSF CCF 2007036. This work was supported by NSF CCF 2312774.

References

- D. Aloise, A. Deshpande, P. Hansen, and P. Popat. Np-hardness of euclidean sum-of-squares clustering. *Machine learning*, 75(2):245–248, 2009.
- G. Aridor, Y. Mansour, A. Slivkins, and Z. S. Wu. Competing bandits: The perils of exploration under competition. *arXiv preprint arXiv:2007.10144*, 2020.
- O. Ben-Porat and M. Tennenholtz. Regression equilibrium. In *Proceedings of the 2019 ACM Conference on Economics and Computation*, pages 173–191, 2019.
- N. Bof, R. Carli, and L. Schenato. Lyapunov theory for discrete time systems. *arXiv preprint arXiv:1809.05289*, 2018.
- A. Bose, M. Curmei, D. L. Jiang, J. Morgenstern, S. Dean, L. J. Ratliff, and M. Fazel. Initializing services in interactive ml systems for diverse users. *arXiv preprint arXiv:2312.11846*, 2023.
- 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.
- 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, 2021.
- F. Ding, M. Hardt, J. Miller, and L. Schmidt. Retiring adult: New datasets for fair machine learning. *Advances in Neural Information Processing Systems*, 34:6478–6490, 2021.
- D. Drusvyatskiy and L. Xiao. Stochastic optimization with decision-dependent distributions. *arXiv preprint arXiv:2011.11173*, 2020.
- D. P. Foster and R. V. Vohra. An economic argument for affirmative action. *Rationality and Society*, 4(2):176–188, 1992.
- T. Ginart, E. Zhang, Y. Kwon, and J. Zou. Competing ai: How does competition feedback affect machine learning? In *International Conference on Artificial Intelligence and Statistics*, pages 1693–1701. PMLR, 2021.
- M. Hardt, N. Megiddo, C. Papadimitriou, and M. Wooters. Strategic classification. In *Proceedings of the 2016 ACM conference on innovations in theoretical computer science*, pages 111–122, 2016.
- M. Hardt, M. Jagadeesan, and C. Mendler-Dünnér. Performative power. *arXiv preprint arXiv:2203.17232*, 2022.
- T. Hashimoto, M. Srivastava, H. Namkoong, and P. Liang. Fairness without demographics in repeated loss minimization. In *International Conference on Machine Learning*, pages 1929–1938. PMLR, 2018.
- L. Hellemo, P. I. Barton, and A. Tomaszewski. Decision-dependent probabilities in stochastic programs with recourse. *Computational Management Science*, 15(3):369–395, 2018.
- H. Hotelling. Stability in competition. *The Economic Journal*, 1929.
- M. Inaba, N. Katoh, and H. Imai. Applications of weighted voronoi diagrams and randomization to variance-based k-clustering. In *Proceedings of the tenth annual symposium on Computational geometry*, pages 332–339, 1994.
- Z. Izzo, L. Ying, and J. Zou. How to learn when data reacts to your model: performative gradient descent. In *International Conference on Machine Learning*, pages 4641–4650. PMLR, 2021.
- M. Kleindessner, P. Awasthi, and J. Morgenstern. Fair k-center clustering for data summarization. In *International Conference on Machine Learning*, pages 3448–3457. PMLR, 2019a.
- M. Kleindessner, S. Samadi, P. Awasthi, and J. Morgenstern. Guarantees for spectral clustering with fairness constraints. In *International Conference on Machine Learning*, pages 3458–3467. PMLR, 2019b.
- H. J. Kushner. Stochastic stability and control. Technical report, Brown Univ Providence RI, 1967.
- J. Liu and Y. Yuan. On almost sure convergence rates of stochastic gradient methods. In *Conference on Learning Theory*, pages 2963–2983. PMLR, 2022.
- C. Mendler-Dünnér, J. Perdomo, T. Zrnic, and M. Hardt. Stochastic optimization for performative prediction. *Advances in Neural Information Processing Systems*, 33: 4929–4939, 2020.
- J. P. Miller, J. C. Perdomo, and T. Zrnic. Outside the echo chamber: Optimizing the performative risk. In *International Conference on Machine Learning*, pages 7710–7720. PMLR, 2021.
- A. Narang, E. Faulkner, D. Drusvyatskiy, M. Fazel, and L. J. Ratliff. Multiplayer performative prediction: Learning in decision-dependent games. *Proceedings of the 24th International Conference on Artificial Intelligence and Statistics (arXiv:2201.03398)*, 2022.
- F. Orabona. Almost sure convergence of sgd on smooth non-convex functions on <https://parameterfree.com/>, 2020.

- S. H. Owen and M. S. Daskin. Strategic facility location: A review. *European journal of operational research*, 111 (3):423–447, 1998.
- J. Perdomo, T. Zrnic, C. Mendler-Dünner, and M. Hardt. Performative prediction. In *International Conference on Machine Learning*, pages 7599–7609. PMLR, 2020.
- G. Piliouras and F.-Y. Yu. Multi-agent performative prediction: From global stability and optimality to chaos. *arXiv preprint arXiv:2201.10483*, 2022.
- J. Quiñero-Candela, M. Sugiyama, A. Schwaighofer, and N. D. Lawrence, editors. *Dataset shift in machine learning*. Mit Press, 2008.
- M. Ray, D. Drusvyatskiy, M. Fazel, and L. J. Ratliff. Decision-dependent risk minimization in geometrically decaying dynamic environments. In *Proceedings of the Association for the Advancement of Artificial Intelligence Conference on AI (AAAI)*, 2022.
- W. H. Sandholm. Evolutionary game theory. *Complex Social and Behavioral Systems: Game Theory and Agent-Based Models*, pages 573–608, 2020.
- S. Z. Selim and M. A. Ismail. K-means-type algorithms: A generalized convergence theorem and characterization of local optimality. *IEEE Transactions on pattern analysis and machine intelligence*, (1):81–87, 1984.
- K. Wood and E. Dall’Anese. Stochastic saddle point problems with decision-dependent distributions. *arXiv preprint arXiv:2201.02313*, 2022.
- K. Wood, G. Bianchin, and E. Dall’Anese. Online projected gradient descent for stochastic optimization with decision-dependent distributions. *IEEE Control Systems Letters*, 2021.
- X. Zhang, M. Khaliligarekani, C. Tekin, and M. Liu. Group retention when using machine learning in sequential decision making: the interplay between user dynamics and fairness. *Advances in Neural Information Processing Systems*, 32, 2019.

Checklist

1. For all models and algorithms presented, check if you include:
 - (a) A clear description of the mathematical setting, assumptions, algorithm, and/or model. [Yes] Assumptions are stated in Section 3 and in theoretical statements.
 - (b) An analysis of the properties and complexity (time, space, sample size) of any algorithm. [Not Applicable] We do not propose an algorithm.
 - (c) (Optional) Anonymized source code, with specification of all dependencies, including external libraries. [Yes] A link is provided in section 5.
2. For any theoretical claim, check if you include:
 - (a) Statements of the full set of assumptions of all theoretical results. [Yes] Assumptions are stated in Section 3 and in theoretical statements.
 - (b) Complete proofs of all theoretical results. [Yes] Proofs are presented in the appendix.
 - (c) Clear explanations of any assumptions. [Yes] Assumptions are discussed in Section 3.
3. For all figures and tables that present empirical results, check if you include:
 - (a) The code, data, and instructions needed to reproduce the main experimental results (either in the supplemental material or as a URL). [Yes] A link is provided in section 5.
 - (b) All the training details (e.g., data splits, hyperparameters, how they were chosen). [Yes] The details are available in the code.
 - (c) A clear definition of the specific measure or statistics and error bars (e.g., with respect to the random seed after running experiments multiple times). [Yes]
 - (d) A description of the computing infrastructure used. (e.g., type of GPUs, internal cluster, or cloud provider). [Yes]
4. If you are using existing assets (e.g., code, data, models) or curating/releasing new assets, check if you include:
 - (a) Citations of the creator If your work uses existing assets. [Yes]
 - (b) The license information of the assets, if applicable. [Not Applicable] The census data we use is in the public domain.
 - (c) New assets either in the supplemental material or as a URL, if applicable. [Not Applicable]
 - (d) Information about consent from data providers/curators. [Not Applicable] The census data we use is in the public domain.
 - (e) Discussion of sensible content if applicable, e.g., personally identifiable information or offensive content. [Not Applicable]
5. If you used crowdsourcing or conducted research with human subjects, check if you include:
 - (a) The full text of instructions given to participants and screenshots. [Not Applicable]
 - (b) Descriptions of potential participant risks, with links to Institutional Review Board (IRB) approvals if applicable. [Not Applicable]
 - (c) The estimated hourly wage paid to participants and the total amount spent on participant compensation. [Not Applicable]

Supplementary Materials: Emergent segmentation from participation dynamics and multi-learner retraining

A Motivating Examples

We discuss several real-world examples which exhibit degrees of market segmentation across characteristics such as nationality, age, and race. In these examples, market conditions are certainly affected by more complex phenomena, from network effects to explicit competition between firms. While we do not claim that the dynamics we study are necessarily the main contributing factor, our simple model isolates the potential contribution of learning dynamics: namely, to reinforce such segmentation. This perspective highlights the potential effects of efforts to incorporate data or improve personalization.

A.1 Social Media

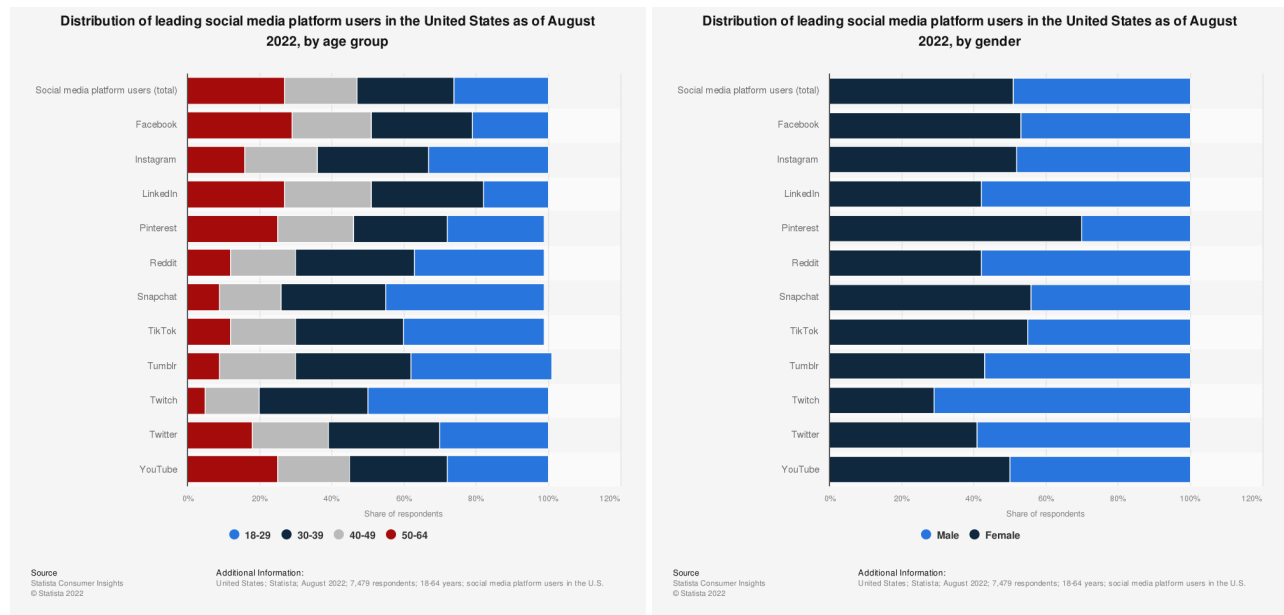


Figure 6: Social media usage across leading social media platforms. Left: Age distribution. Right: Gender distribution

Usage of various social media sites in the US varies across genders³ and age groups⁴. For example, the users of Facebook and LinkedIn skew older while Snapchat, TikTok, Tumblr, and Twitch are more heavily used by the younger population. Similarly users of Pinterest strongly skew female while users of Twitch are more likely to be male. Figure 6 shows the disparities along gender and age for leading social media platforms. These disparities across platforms are reinforced by user behaviors: imagining the experience of a 45 year old logging onto Twitch for the first time compared with a 14 year old; or instead imagine a 14 year old logging into Facebook. Because the usage patterns determine the data available to the platforms, the disparities are also reinforced by the behavior of the platforms themselves. Similarly, Pinterest algorithms are more likely to be tailored to the tastes of a female demographic, while Twitch's to a demographic more representative of males.

³<https://www.statista.com/statistics/1337563/us-distribution-leading-social-media-platforms-by-gender/>

⁴<https://www.statista.com/statistics/1337525/us-distribution-leading-social-media-platforms-by-age-group/>

A.2 Music Streaming

Worldwide market share of music streaming services is split between several companies (see Figure 7). However, the distribution of music streaming by country shows clear patterns: most users in China use Tencent, most users in Mexico use Spotify, and most users in the Middle East and Northern Africa (MENA) use Anghami. On the other hand, the markets United States, Russia, and India are not dominated by a single service. However, the handful of most used services in these regions have a small market outside of their main market. Due to this segmented market, only certain platforms collect large scale data about music preferences in certain regions. If many users from western cultures make playlists containing both Arabic and Indian music, Spotify may learn to associate those genres in a way that is undesirable or even offensive to users from those cultures. This leads to a self-reinforcing effect: services who make bad predictions for users from certain cultures are unlikely to correct this bias as those users choose instead to use services that more accurately reflect their tastes.

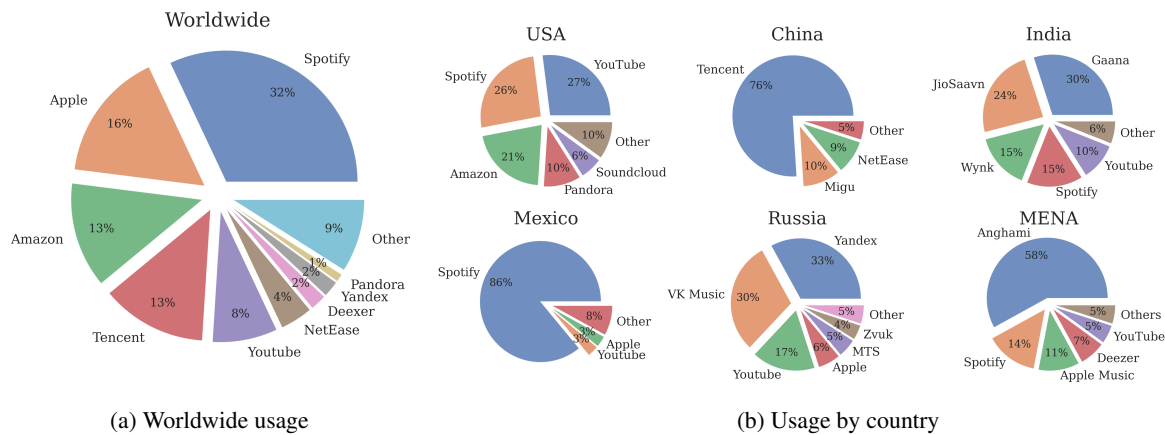


Figure 7: Usage of music streaming services in different markets⁵. Left: Worldwide market share. Right: Market share in USA, China, India, Mexico, Russia, Middle East and Northern Africa (MENA).

A.3 Personalized health

The growing popularity of direct-to-consumer genetic testing is driven by the growth of two market leaders: AncestryDNA and 23andMe⁶. These tests are used both for determining ancestry as well as receiving polygenic risk scores for various medical conditions. The accuracy of the tests varies across ethnic groups; with Latino, Middle Eastern and, African ancestry being most under-represented. This issue is self re-inforcing; for instance people of African descent are less likely to use a large service like 23andMe and more likely to use a specialized service such as AfricanAncestry⁷.

⁵ All statistics recorded from Statista:

Worldwide: <https://www.statista.com/statistics/653926/music-streaming-service-subscriber-share/>,

United States: <https://www.statista.com/statistics/1351506/streaming-services-music-podcasts-united-states/>,

China: <https://www.statista.com/statistics/711295/china-leading-mobile-music-platforms-by-active-user-number/>,

India: <https://www.statista.com/statistics/922400/india-music-app-market-share/>,

Mexico: <https://www.statista.com/statistics/1018370/over-the-top-audio-platforms-mexico-by-market-share/>,

Russia: <https://www.statista.com/statistics/1347035/most-popular-music-streaming-platforms-in-russia/>,

Middle East and North America:

<https://www.statista.com/statistics/1295716/mena-share-of-paying-music-streaming-subscribers-by-platform/>,

⁶ <https://www.statista.com/chart/17023/commercial-genetic-testing/>

⁷ <https://africanancestry.com/>

B Preliminaries

B.1 Notation

We introduce a compact notation. The simplex product is defined as

$$\Delta_m^n = \left\{ A \in \mathbb{R}^{n \times m} \mid \sum_{j=1}^m A_{ij} = 1 \right\}$$

so that the rows sum to 1. Then the state space of subpopulation allocations and learner parameters is $\mathcal{X} = \Delta_m^n \times \mathbb{R}^{m \times d}$. For a square matrix A , we use the notation $\text{diag}(A)$ to represent the vector containing the diagonal entries of A . For a vector a , $\text{Diag}(a)$ is a diagonal matrix with a along the diagonal. Furthermore we will say $a \leq b$ for vectors a, b if the inequality holds elementwise.

Define a matrix valued risk function $R : \mathbb{R}^{m \times d} \rightarrow \mathbb{R}^{n \times m}$ so that $R_{ij}(\Theta) = \mathcal{R}_i(\theta_j)$. Recall that in Section 3.1 the subpopulation and learner risks played a key role. We therefore define vector valued functions $\bar{\mathcal{R}}^{\text{subpop}} : \mathcal{X} \rightarrow \mathbb{R}^n$ and $\bar{\mathcal{R}}^{\text{learner}} : \mathcal{X} \rightarrow \mathbb{R}^m$ as follows:

$$\bar{\mathcal{R}}^{\text{subpop}}(\alpha, \Theta) = \text{diag}(\alpha R(\Theta)^\top), \quad \bar{\mathcal{R}}^{\text{learner}}(\alpha, \Theta) = \text{diag}(\text{Diag}(\alpha^\top \beta)^{-1} \alpha^\top \text{Diag}(\beta) R(\Theta)).$$

Then the definition of risk reducing dynamics for subpopulations and learners can be written as

$$\bar{\mathcal{R}}^{\text{subpop}}(\alpha^{t+1}, \Theta) \leq \bar{\mathcal{R}}^{\text{subpop}}(\alpha^t, \Theta) \quad \text{and} \quad \bar{\mathcal{R}}^{\text{learner}}(\alpha, \Theta^{t+1}) \leq \bar{\mathcal{R}}^{\text{learner}}(\alpha, \Theta^t).$$

Risk minimizing in the limit is defined similarly, where the inequality is strict for at least one entry of the vectors unless the state is at a local minimum.

The total risk can be written as

$$\mathcal{R}^{\text{total}}(\alpha, \Theta) := \text{tr}(\text{diag}(\beta) \alpha R(\Theta)^\top).$$

Lemma B.1. *Under the assumption that all loss functions are continuous, the risk function R is continuous w.r.t. to Θ , and thus $\mathcal{R}^{\text{total}}$ is continuous w.r.t. α and Θ .*

The sequential dynamics updates described in Section 3.1 can be written as

$$\begin{bmatrix} \alpha^{t+1} \\ \Theta^{t+1} \end{bmatrix} = \begin{bmatrix} \nu(\alpha^t, \Theta^t) \\ \mu(\alpha^{t+1}, \Theta^t) \end{bmatrix} = \begin{bmatrix} \nu(\alpha^t, \Theta^t) \\ \mu(\nu(\alpha^t, \Theta^t), \Theta^t) \end{bmatrix} =: f(\alpha^t, \Theta^t). \quad (4)$$

Lemma B.2. *As long as the subpopulation and learner updates described in Section 3.1 are locally Lipschitz, so is the dynamics function f defined in (4).*

B.2 Background

For completeness, we include important results and definitions that our proofs will make use of. First, we state two theorems about Lyapunov theory for stability.

Theorem B.3 (Theorem 1.2 in Bof et al. (2018)). *Let $x_{\text{eq}} \in \mathcal{D}$ be an equilibrium point for the autonomous systems $x_{t+1} = f(x_t)$ where $f : \mathcal{D} \rightarrow \mathcal{X}$ is locally Lipschitz in $\mathcal{D} \subseteq \mathcal{X}$. Suppose there exists a function $V : \mathcal{D} \rightarrow \mathbb{R}$ which is continuous and such that*

$$\begin{aligned} V(x_{\text{eq}}) &= 0 \quad \text{and} \quad V(x) > 0 \quad \forall \quad x \in \mathcal{D} - \{x_{\text{eq}}\} \\ V(f(x)) - V(x) &\leq 0 \quad (\text{resp.} < 0) \quad \forall \quad x \in \mathcal{D} \end{aligned}$$

Then x_{eq} is stable (resp. asymptotically stable).

Theorem B.4 (Theorem 1.5 in Bof et al. (2018)). *Let $x_{\text{eq}} \in \mathcal{D}$ be an equilibrium point for the autonomous systems $x_{t+1} = f(x_t)$ where $f : \mathcal{D} \rightarrow \mathcal{X}$ is locally Lipschitz in $\mathcal{D} \subseteq \mathcal{X}$. Let $V : \mathcal{D} \rightarrow \mathbb{R}$ be a continuous function with $V(x_{\text{eq}}) = 0$ and $V(x_0) > 0$ for some x_0 arbitrarily close to x_{eq} . Let $r > 0$ be such that $B_r(x_{\text{eq}}) \subseteq \mathcal{D}$ and $\mathcal{U} = \{x \in B_r(x_{\text{eq}}) \mid V(x) > 0\}$, and suppose that $V(f(x)) - V(x) > 0$ for all $x \in \mathcal{U}$. Then x_{eq} is not stable.*

Next, we state the definition of a (isolated) local minimum.

Definition B.5. The point u_* is a local minimum (resp. isolated local minimum) of a function h over a domain $\mathcal{U}$ if there is a $\delta > 0$ such that for any $u \in \mathcal{U}$ with $\|u - u_*\| \leq \delta$, $h(u_*) \leq h(u)$ (resp. $h(u_*) < h(u)$).

Next, we state the implicit function theorem.

Theorem B.6 (Implicit Function Theorem). *Let $U \subseteq \mathbb{R}^n$, $V \subseteq \mathbb{R}^m$ be open sets and $f : U \times V \rightarrow \mathbb{R}$ is C^r for some $r \geq 1$. For some $x_0 \in U$, $y_0 \in V$ assume the partial derivative in the second argument $D_2 f(x_0, y_0) : \mathbb{R}^m \rightarrow \mathbb{R}$ is an isomorphism. Then there are neighborhoods U_0 of x_0 and W_0 of $f(x_0, y_0)$ and a unique C^r map $g : U_0 \times W_0 \rightarrow V$ such that for all $(x, w) \in U_0 \times W_0$, $f(x, g(x, w)) = w$.*

Finally we prove a property of intersecting convex hulls.

Lemma B.7. *Let $x_1, x_2, \dots, x_n$ and $y_1, y_2, \dots, y_m$ be some points in $\mathbb{R}^d$. Define by $\mathcal{C}_x$ and $\mathcal{C}_y$ the convex hulls of $\{x_i\}_{i=1}^n$ and $\{y_j\}_{j=1}^m$ respectively. Then there do not exist points $\bar{x} \in \mathbb{R}^d$ and $\bar{y} \in \mathbb{R}^d$ such that the following inequalities are satisfied:*

$$\begin{aligned} \|x_i - \bar{x}\| &< \|x_i - \bar{y}\| \quad \forall i = 1, 2, \dots, n \\ \|y_j - \bar{y}\| &< \|y_j - \bar{x}\| \quad \forall j = 1, 2, \dots, m \end{aligned}$$

Proof. Assume by contradiction that the inequalities above hold. Define $\mathcal{H}_x := \{z \in \mathbb{R}^d \mid \|z - \bar{x}\| < \|z - \bar{y}\|\}$ and $\mathcal{H}_y := \{z \in \mathbb{R}^d \mid \|z - \bar{y}\| < \|z - \bar{x}\|\}$. The sets $\mathcal{H}_x$ and $\mathcal{H}_y$ are disjoint half-spaces (without boundary) then defined by the hyperplane bisecting the segment connecting $\bar{x}$ and $\bar{y}$. By assumption then we have that $x_i \in \mathcal{H}_x$ for all i and $y_j \in \mathcal{H}_y$ for all j ; since $\mathcal{H}_x$ and $\mathcal{H}_y$ are convex, it follows that $\mathcal{C}_x \subset \mathcal{H}_x$ and $\mathcal{C}_y \subset \mathcal{H}_y$. Therefore $\mathcal{C}_x \cap \mathcal{C}_y = \emptyset$, which leads to a contradiction. $\square$

B.3 Properties of partial minimization

In this section, we state a handful of important results about the partial minimization of the total risk. This is somewhat similar to the analysis presented by [Selim and Ismail \(1984\)](#) in the context of clustering algorithms.

Lemma B.8. *Define the function $F : \mathbb{R}^{m \times n} \rightarrow \mathbb{R}$ as $F(\alpha) = \min_{\Theta} \mathcal{R}^{\text{total}}(\alpha, \Theta)$. This function is concave and a point (α^0, Θ^0) is a local minimum of $\mathcal{R}^{\text{total}}$ over the domain $\mathcal{X} = \mathcal{X}_\alpha \times \mathbb{R}^{m \times d}$ if α^0 is a local minimum of F over the domain $\mathcal{X}_\alpha$ and $\Theta^0 \in \arg \min_{\Theta} \mathcal{R}^{\text{total}}(\alpha^0, \Theta)$. Furthermore, in the case that Θ^0 is the unique minimizer of $\mathcal{R}^{\text{total}}(\alpha^0, \Theta)$, then (α^0, Θ^0) is a local minimum (resp. isolated local minimum) if and only if α^0 is a local minimum (resp. isolated local minimum).*

Proof of Lemma B.8. $F(\alpha)$ is well defined due to the convexity of the risk functions. Concavity follows from the observation that F is the point-wise minimum of a family of functions which are linear in α (since for every fixed Θ , the total risk is linear in α).

We break the proof of equivalence into two implications.

1. F minimized $\implies \mathcal{R}^{\text{total}}$ minimized

There is a $\delta > 0$ such that for any $\alpha \in \mathcal{X}_\alpha$ with $\|\alpha^0 - \alpha\| \leq \delta$, $F(\alpha^0) \leq F(\alpha)$, i.e.

$$\mathcal{R}^{\text{total}}(\alpha^0, \Theta^0) \leq \mathcal{R}^{\text{total}}(\alpha, \Theta^*(\alpha))$$

for any minimizing $\Theta^*(\alpha)$. For fixed allocation α define $\mathcal{R}_\alpha^{\text{total}}(\Theta) = \mathcal{R}^{\text{total}}(\alpha, \Theta)$ which is convex and minimized at $\Theta^*(\alpha)$ and hence:

$$\mathcal{R}^{\text{total}}(\alpha, \Theta^*(\alpha)) \leq \mathcal{R}^{\text{total}}(\alpha, \Theta), \quad \forall \Theta.$$

Combining the inequalities yields: $\mathcal{R}^{\text{total}}(\alpha^0, \Theta^0) \leq \mathcal{R}^{\text{total}}(\alpha, \Theta)$, and thus (α^0, Θ^0) is a local minimum of $\mathcal{R}^{\text{total}}$. The implication for the isolated local minimum case follows by the same arguments with strict inequalities on the total risk, noting that if Θ^0 is a unique minimizer, it must also be isolated.

2. $\mathcal{R}^{\text{total}}$ minimized $\implies F$ minimized

Recall that $\mathcal{R}^{\text{total}}(\alpha, \Theta)$ can be written as $\text{tr}(\text{diag}(\beta)\alpha R(\Theta)^\top)$. Then

$$\mathcal{R}^{\text{total}}(\alpha^0 + D, \Theta^0) - \mathcal{R}^{\text{total}}(\alpha^0, \Theta^0) = \text{tr}(\text{diag}(\beta)DR(\Theta^0)^\top) \geq 0$$

where inequality holds for all $D \in \mathbb{R}^{n \times m}$ such that $\alpha^0 + D \in \mathcal{X}_\alpha$ by the fact that α^0 is a minimum. Recognizing the gradient from Lemma B.9 and using the uniqueness of Θ^0 , the expression is equivalently $\langle \nabla_\alpha F(\alpha^0), D \rangle \geq 0$. In other words, the directional derivative in any feasible direction D is non-negative. Hence, α^0 is a local minimum of F . The implication for the isolated local minimum case follows by the same arguments with strict inequalities on the total risk. $\square$

Lemma B.9. For $F : \mathbb{R}^{n \times m} \rightarrow \mathbb{R}$ defined as in Lemma B.8 suppose the minimizer $\Theta^*(\alpha) = \arg \min_\Theta \mathcal{R}^{\text{total}}(\alpha, \Theta)$ is unique. The gradient is

$$\nabla_\alpha F(\alpha) = \text{diag}(\beta) R(\Theta^*(\alpha)), \text{ i.e. } \frac{\partial F(\alpha)}{\partial \alpha_{ij}} = \beta_i \mathcal{R}_i(\theta_j^*(\alpha)).$$

Further suppose that the risks are strongly convex. Then second partial derivatives are given by

$$\frac{\partial^2 F(\alpha)}{\partial \alpha_{k\ell} \partial \alpha_{ij}} = \begin{cases} 0 & k \neq j \\ -\beta_i \nabla \mathcal{R}_i(\theta_j^*)^\top (\sum_{\ell'} \beta_{\ell'} \alpha_{\ell'j} \nabla^2 \mathcal{R}_{\ell'}(\theta_j^*)) \nabla \mathcal{R}_\ell(\theta_j^*) & k = j \end{cases}.$$

Proof. Computing the gradient:

$$\nabla_\alpha F(\alpha) = \nabla_\alpha \mathcal{R}^{\text{total}}(\alpha, \Theta^*(\alpha)) + \nabla \Theta^*(\alpha) \nabla_\theta \mathcal{R}^{\text{total}}(\alpha, \Theta^*(\alpha)) = \text{diag}(\beta) R(\Theta).$$

The first equality follows by product rule. The second equality follows because 1) the total risk is linear in α and 2) the second term is zero due to the optimality of $\Theta^*(\alpha)$.

Now notice that

$$\frac{\partial}{\partial \alpha_{k\ell}} \mathcal{R}_i(\theta_j^*(\alpha)) = \left\langle \frac{\partial \theta_j^*(\alpha)}{\partial \alpha_{k\ell}}, \nabla_\theta \mathcal{R}_i(\theta_j^*(\alpha)) \right\rangle$$

To compute the derivatives of $\theta_j^*(\alpha)$ we use the implicit function theorem and the assumption that the risks are strongly convex. We apply the implicit function theorem to the first order optimality condition

$$\theta_j^*(\alpha) \in \arg \min_{\theta_j} \bar{\mathcal{R}}_j^{\text{learner}}(\alpha_{:,j}, \theta_j)$$

The Hessian $\nabla_\theta^2 \bar{\mathcal{R}}_j^{\text{learner}}(\alpha, \Theta)$ is non-degenerate due to strong convexity of the subpopulation risks. There exists a neighborhood U_0 of α and a unique (sufficiently smooth) map $\theta_j^*(\cdot)$ such that for all $\alpha \in U_0$, we have that $\nabla_\theta \bar{\mathcal{R}}_j^{\text{learner}}(\alpha, \theta_j^*(\alpha)) = 0$. Then by implicit function theorem we obtain

$$\nabla \theta_j^*(\alpha) = -\nabla_\theta^2 \bar{\mathcal{R}}_j^{\text{learner}} \circ \nabla_{\alpha\theta} \bar{\mathcal{R}}_j^{\text{learner}}(\alpha_{:,j}, \theta_j^*(\alpha))$$

by taking the derivative of the first order condition differentiating through $\theta_j^*(\cdot)$ and setting it to zero. We have that

$$\nabla_\theta^2 \bar{\mathcal{R}}_j^{\text{learner}} = \sum_{\ell'} \beta_{\ell'} \alpha_{\ell'j} \nabla^2 \mathcal{R}_{\ell'}(\theta_j^*), \quad \frac{\partial}{\partial \alpha_{k\ell}} \nabla_\theta \bar{\mathcal{R}}_j^{\text{learner}} = \begin{cases} 0 & k \neq j \\ \nabla \mathcal{R}_\ell(\theta_j^{\text{eq}}) & k = j \end{cases}.$$

The result follows by combining the expressions. $\square$

C Full Proofs of Main Results

In this section, we present proofs of the main results.

C.1 Connections between dynamics and total risk

Proposition 4.2. For any risk-reducing subpopulation and learner dynamics, the total risk is non-increasing: $\mathcal{R}^{\text{total}}(\alpha^{t+1}, \Theta^{t+1}) \leq \mathcal{R}^{\text{total}}(\alpha^t, \Theta^t), \forall t$. If subpopulations and learners are risk minimizing in the limit, then the total risk is strictly decreasing unless (α^t, Θ^t) is a local minimizer of $\mathcal{R}^{\text{total}}$.

Proof of Proposition 4.2 The key to seeing that the total risk acts like a potential for the market dynamics is to note two equivalent decompositions of the total risk:

$$\mathcal{R}^{\text{total}}(\alpha, \Theta) = \beta^\top \bar{\mathcal{R}}^{\text{subpop}}(\alpha, \Theta) = \beta^\top \alpha \bar{\mathcal{R}}^{\text{learner}}(\alpha, \Theta).$$

Being risk-reducing learners' updates satisfy:

$$\bar{\mathcal{R}}^{\text{learner}}(\alpha^t, \Theta^{t+1}) \leq \bar{\mathcal{R}}^{\text{learner}}(\alpha^t, \Theta^t) \implies \mathcal{R}^{\text{total}}(\alpha^t, \Theta^{t+1}) \leq \mathcal{R}^{\text{total}}(\alpha^t, \Theta^t).$$

Similarly risk reducing subpopulations satisfy:

$$\bar{\mathcal{R}}^{\text{subpop}}(\alpha^{t+1}, \Theta^{t+1}) \leq \bar{\mathcal{R}}^{\text{subpop}}(\alpha^t, \Theta^{t+1}) \implies \mathcal{R}^{\text{total}}(\alpha^{t+1}, \Theta^{t+1}) \leq \mathcal{R}^{\text{total}}(\alpha^t, \Theta^{t+1}).$$

Finally, combining the two updates yields the desired inequality.

In the case that learners and subpopulations are risk minimizing in the limit, the same argument holds with strict inequality, unless $(\alpha^{t+1}, \Theta^{t+1})$ is a local minimum. $\square$

Theorem 4.3. *For any learners and subpopulations who are risk minimizing in the limit, an equilibrium $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is asymptotically stable if it is an isolated local minimizer of the total risk $\mathcal{R}^{\text{total}}$. If it is not a local minimizer of the total risk, then it is not stable.*

Proof of Theorem 4.3 We break this proof into two implications.

1. Isolated local min $\implies$ Asymptotic stability

Define $V(\alpha, \Theta) = \mathcal{R}^{\text{total}}(\alpha, \Theta) - \mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta^{\text{eq}})$. The dynamics f are Lipschitz by Lemma B.2 and this V satisfies the conditions of Theorem B.3 with strict inequality, thus we conclude that $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is an asymptotically stable equilibrium.

2. Not local min $\implies$ Not stable

Define $V(\alpha, \Theta) = \mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta^{\text{eq}}) - \mathcal{R}^{\text{total}}(\alpha, \Theta)$ which will increase along trajectories. Since we are not at a local min, there must be some arbitrarily close α^0, Θ^0 such that $V(\alpha, \Theta) > 0$. Then we apply Theorem B.4 which guarantees that the equilibrium is not stable. $\square$

Corollary 4.4. *Equilibria exist when learners and subpopulations are risk minimizing in the limit and the total risk function has isolated local minima. They may not exist otherwise.*

Proof of Corollary 4.3 We first argue that if the dynamics are risk minimizing, then an isolated local minimum of the total risk must be an equilibria. Let (α^0, Θ^0) denote the isolated local minima of the total risk. It must be that α^0 is an isolated, and thus unique, minimizer of $\mathcal{R}^{\text{total}}(\alpha, \Theta^0)$ since the function is linear in α . We can thus conclude that $\nu(\alpha^0, \Theta^0) = \alpha^0$. It also must be that Θ^0 is a unique minimizer of $\mathcal{R}^{\text{total}}(\alpha^0, \Theta)$ since the function is convex in Θ . We can thus conclude that $\mu(\alpha^0, \Theta^0) = \Theta^0$. Therefore (α^0, Θ^0) is equilibrium of the dynamics.

We next show that equilibria may not exist when the dynamics are not risk minimizing in the limit. To show that they may not exist otherwise, consider the following example. Let all learners be static and identical so $\Theta^{t+1} = \Theta^t$ and $\Theta = (\theta, \theta, \dots, \theta)$. Let the subpopulation update break ties among equivalent learners randomly. Then the subpopulations will randomly switch between learners. Though these dynamics satisfy the definition of risk reducing (at equality), they will not converge to an equilibrium.

We lastly show that equilibria may not exist when the total risk function does not have an isolated local minima. Suppose that learners update with full risk minimization and all subpopulations have risk uniquely minimized at the same value θ . Finally suppose that subpopulations will break ties among equivalent learners randomly (and are otherwise risk minimizing). As in the previous example, the subpopulations will randomly switch between learners and no equilibrium exists. $\square$

C.2 Stable equilibria

Theorem 4.6. *Suppose learners and subpopulations are risk minimizing in the limit, α^{eq} is segmented, and $\mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta)$ has a unique minimizer Θ^{eq} . Define a mapping $\gamma : [n] \rightarrow [m]$ such that $\gamma(i) = j$ is the learner with nonzero mass in $\alpha_{i,j}^{\text{eq}}$. If every subpopulation strictly prefers their current learner:*

$$\mathcal{R}_i(\theta_{\gamma(i)}^{\text{eq}}) < \mathcal{R}_i(\theta_j^{\text{eq}}), \quad (1)$$

for all i and learners $j \neq \gamma(i)$, then $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is an asymptotically stable equilibrium. If there is a subpopulation who would strictly prefer to switch learners, then $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is not stable.

Proof of Theorem 4.6 First note that it must be that every learner is associated to at least one subpopulation. Otherwise, the total risk would not have a unique minimizer over Θ .

We start with the first statement, and show that the stated conditions imply that $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ is isolated local minimum of the total risk. By Theorem 4.3, this implies asymptotic stability.

We specifically argue the conditions are sufficient for guaranteeing an isolated local minimum with respect to $F(\alpha)$, appealing to Lemma B.8. First notice that we have the unique $\Theta^{\text{eq}} = \arg \min_{\Theta} \mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta)$ as required. Suppose by contradiction that there is some perturbation to α that causes $F(\alpha)$ to decrease or remain the same. Equivalently, the projection of the negative gradient onto the simplex points towards some other vertex, i.e. the component of the gradient in the direction of learner j is less than or equal to that in the direction of $\gamma(i)$ for some $j \neq \gamma(i)$. We can write this condition as

$$\frac{\partial F(\alpha)}{\partial \alpha_{i\gamma(i)}} \geq \frac{\partial F(\alpha)}{\partial \alpha_{ij}} \iff \mathcal{R}_i(\theta_{\gamma(i)}^{\text{eq}}) \geq \mathcal{R}_i(\theta_j^{\text{eq}})$$

where we use Lemma B.9. This violates the risk comparison condition (I), and therefore there must be no such perturbation, and thus α^{eq} is an isolated local minimum.

We turn our attention to the second statement. Theorem 4.3, it is equivalent to argue about minima of the total risk function. Suppose that for some subpopulation, there is some learner for which $\mathcal{R}_i(\theta_{\gamma(i)}^{\text{eq}}) > \mathcal{R}_i(\theta_j^{\text{eq}})$. Then any small perturbation of that subpopulations's allocation towards that learner will decrease the total risk, and thus the point is not a minimum. $\square$

In a segmented allocation, each θ_j^{eq} will minimize the average loss over the group of subpopulations assigned to them. Denote the parameter which minimizes risk of subpopulation i as $\phi_i := \arg \min_{\theta \in \mathbb{R}^d} \mathcal{R}_i(\theta)$. Then each θ_j^{eq} is a convex combination of ϕ_i for i in j th partition. Using this perspective, we provide an intuitive necessary (but not sufficient) condition for a class of symmetric risk functions.

Corollary C.1. Suppose that risk functions satisfy $\mathcal{R}_i(\theta) < \mathcal{R}_i(\theta') \iff \|\theta - \phi_i\| < \|\theta' - \phi_i\|$ for ϕ_i the subpopulation optimal parameter. Then in an asymptotically stable segmented equilibrium, the convex hulls of the grouped subpopulations optimal parameters $\{\phi_i\}$ are non-intersecting.

Applying this Corollary to the example in Figure 2, we see that a segmented equilibrium with subpopulation 1 and 3 participating in the same learner cannot be stable.

Proof of Corollary C.1. Let $\phi_1, \phi_2, \dots, \phi_k \in \mathbb{R}^d$ be the optimal decisions for the subpopulations allocated to the first learner and $\psi_1, \psi_2, \dots, \psi_l \in \mathbb{R}^d$ be the optimal decisions for the subpopulations allocated to the second learner. Let θ_1 and θ_2 be the decisions of each learner. Assume that the convex hulls of $\{\phi_i\}_{i=1}^k$ and $\{\psi_i\}_{i=1}^l$ intersect. By Lemma B.7, there exists i such that $\|\phi_i - \theta_2\| \leq \|\phi_i - \theta_1\|$. By the assumption about the risk functions, this implies $\mathcal{R}_i(\theta_2) < \mathcal{R}_i(\theta_1)$. In other words, there exist a subpopulation that would prefer to switch learners. Thus by Theorem 4.6 these allocation of subpopulations to learner is not stable and so the convex hulls must not intersect. $\square$

Theorem 4.8. Consider dynamics which are risk minimizing in the limit and an α^{eq} with any subpopulation i having nonzero support on set of two or more learners $j \in \mathcal{J}$. Assume risks are strongly convex and define $\Theta^{\text{eq}} = \arg \min \mathcal{R}^{\text{total}}(\alpha^{\text{eq}}, \Theta)$. Then $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$ cannot be stable unless it is “balanced” in the sense that learners in $\mathcal{J}$ are risk equivalent and optimal for i , i.e. for all $j, j' \in \mathcal{J}$,

$$\mathcal{R}_i(\theta_j^{\text{eq}}) = \mathcal{R}_i(\theta_{j'}^{\text{eq}}) \quad \text{and} \quad \nabla \mathcal{R}_i(\theta_j^{\text{eq}}) = 0. \quad (2)$$

If it is balanced, so are all allocations for subpopulation i with support over $\mathcal{J}$. Finally, all stable equilibria must be either balanced or segmented.

Proof of Theorem 4.8 Theorem 4.3 shows that an equilibrium cannot be stable if it is not a local minimum of the total risk. We therefore develop conditions under which an equilibrium point will be a local minimum. By Lemma B.8, it is equivalent to argue about the local minima of the concave function $F(\alpha)$ over the simplex product Δ_m^n . All minima of the total risk will occur on the boundary of the simplex product, i.e. a face or a vertex. Since F is still concave when restricted to a face of the simplex, the same argument shows the minima are on the boundary, hence vertices, except for the degenerate case where F takes a constant value over the face.

We now characterize this degenerate case. F takes a constant value over the face if and only if 1) the gradient of F is perpendicular to the face at α and 2) remains perpendicular along the face. The face is described by a set of indices $\mathcal{J} \subseteq [m]$. Mathematically, we write the two conditions as: for all pairs $j, j' \in \mathcal{J}$, $\ell \in [n]$, and $k \in [m]$,

$$\frac{\partial F(\alpha)}{\partial \alpha_{ij}} = \frac{\partial F(\alpha)}{\partial \alpha_{ij'}} \quad \text{and} \quad \frac{\partial}{\partial \alpha_{\ell k}} \left(\frac{\partial F(\alpha)}{\partial \alpha_{ij}} - \frac{\partial F(\alpha)}{\partial \alpha_{ij'}} \right) = 0 \quad (5)$$

Using Lemma B.8, the first expression simplifies to the *risk equivalent* condition that $\mathcal{R}_i(\theta_j^{\text{eq}}) = \mathcal{R}_i(\theta_{j'}^{\text{eq}})$. Turning to the second expression in (5), we first use Lemma B.9 to compute

$$\frac{\partial}{\partial \alpha_{\ell k}} \frac{\partial F(\alpha)}{\partial \alpha_{ij}} = \begin{cases} 0 & k \neq j \\ -\beta_i \nabla \mathcal{R}_i(\theta_j^{\text{eq}})^\top \left(\sum_{\ell'} \beta_{\ell'} \alpha_{\ell' j} \nabla^2 \mathcal{R}_{\ell'}(\theta_j^{\text{eq}}) \right) \nabla \mathcal{R}_\ell(\theta_j^{\text{eq}}) & k = j \end{cases}$$

Thus, the condition trivially holds for $k \notin \{j, j'\}$. Otherwise, when $\ell = i$, the condition in (5) requires that

$$\nabla \mathcal{R}_i(\theta_k^{\text{eq}})^\top \left(\sum_{\ell'} \beta_{\ell'} \alpha_{\ell' k} \nabla^2 \mathcal{R}_{\ell'}(\theta_k^{\text{eq}}) \right) \nabla \mathcal{R}_i(\theta_k^{\text{eq}}) = 0, \quad k \in \{j, j'\}$$

Due to the strong convexity of the risks, the Hessian matrix is positive definite. Thus it must be that $\nabla \mathcal{R}_i(\theta_j^{\text{eq}}) = 0$ for all $j \in \mathcal{J}$, i.e. the *risk optimal* condition. Risk optimality implies that the condition holds also when $\ell \neq i$ and thus the characterization is complete. $\square$

C.3 Social Welfare

Proof of Proposition 4.10 Social welfare is non-decreasing (or increasing) if and only if total risk is non-increasing (or decreasing), as guaranteed by Proposition 4.2. Maxima of the social welfare are equivalent to minima of the total risk and therefore the connections to stable equilibria follow by Theorem 4.3. $\square$

Proof of Proposition 4.12 By construction $(\tilde{\alpha}^{\text{eq}}, \tilde{\Theta}^{\text{eq}})$ is not segmented, and neither is it a stable balanced equilibrium (by the non-optimality assumption). Therefore, it is not stable (Theorem 4.8), and thus not a local minimum of the total risk (Theorem 4.3). A perturbation will thus send the system along a risk-reducing trajectory. $\square$

D Detailed Examples

D.1 Risk Reducing and Minimizing Dynamics

Proposition 3.4. *A subpopulation i updating their participation with multiplicative weights is risk minimizing in the limit if $\gamma > 0$ and $\alpha_{ij}^0 > 0 \forall j$. A learner updating its parameter with gradient descent is risk minimizing in the limit when the risk functions $\mathcal{R}_i(\theta)$ are L smooth and the step size satisfies $\gamma^t < \frac{2}{L}$, $\sum_{t=0}^{\infty} \gamma^t = \infty$, and $\sum_{t=1}^{\infty} (\gamma^t)^2 < \infty$.*

Proof. To see that the subpopulation is risk minimizing, first see that

$$\begin{aligned} \bar{\mathcal{R}}_i^{\text{subpop}}(\alpha_{i,:}^{t+1}, \Theta) &= \sum_{j=1}^m \alpha_{ij}^{t+1} \mathcal{R}_i(\theta_j) \\ &= \sum_{j=1}^m \frac{\alpha_{ij}^t \cdot \exp(-\gamma \mathcal{R}_i(\theta_j))}{\sum_{j=1}^m \alpha_{ij}^t \cdot \exp(-\gamma \mathcal{R}_i(\theta_j))} \mathcal{R}_i(\theta_j) \\ &< \sum_{j=1}^m \alpha_{ij}^t \mathcal{R}_i(\theta_j) = \bar{\mathcal{R}}_i^{\text{subpop}}(\alpha_{i,:}^t, \Theta) \end{aligned}$$

where the strict inequality holds as long as α_{ij}^t is not on the boundary of the simplex. Second, observe that for a fixed Θ , $\alpha_{ij}^t \rightarrow 1$ if and only if $\mathcal{R}_i(\theta_j)$ is minimal over all learners for which $\alpha_{ij}^0 > 0$.

To see that the learner is risk minimizing, notice that the gradient update is equivalently

$$\theta_j^{t+1} = \theta_j^t - \gamma_t \nabla_{\theta} \bar{\mathcal{R}}_j^{\text{learner}}(\alpha_{:,j}, \theta_j) .$$

Gradient descent on an L -smooth and convex function leads to strictly decreasing objective values when θ_j^t is not at a minimum and the step size satisfies $\gamma^t < \frac{2}{L}$. It further converges to a minimum in the limit as long as the step size satisfies the Robbins-Munroe condition (see, e.g. [Liu and Yuan \(2022\)](#); [Orabona \(2020\)](#)). $\square$

Example D.1 (Non-continuity of allocation updates). Suppose a population prefers one learner over others, and only shifts participation away from the preferred learner if there is another with risk smaller by at least $R_0 > 0$. This is risk reducing but not minimizing in the limit.

Example D.2 (Shifting to lower-risk models). If a subpopulation's allocation updates always shift allocation from learners with high subpopulation risk to learners with lower subpopulation risk, then the allocation is risk reducing. It may or may not be risk minimizing in the limit.

Example D.3 (Allocations determined by gradient descent). Consider an allocation determined by (projected) gradient descent with respect to a subpopulation's average risk. This is risk-reducing, and may be risk minimizing in the limit depending on the step-size.

D.2 Stability

To illustrate the subtleties of determining stability when the total risk function has non-isolated local minima, we consider a setting with $n = m = 2$ subpopulations and learners where $\mathcal{R}_1(\theta) = \mathcal{R}_2(\theta) = \theta^2$. Then the total risk function is minimized for any $\alpha \in \Delta_m^n$ and $\Theta = (0, 0)$. This continuum of minima can contain equilibria of risk minimizing dynamics, and those equilibria may be stable, asymptotically stable, or unstable, which we illustrate with the following examples.

Example D.4 (Continuum of stable balanced markets). Suppose that subpopulations update their allocation via any Lipschitz continuous risk minimizing update rule which is stationary whenever learners are risk equivalent (i.e. $\mathcal{R}_i(\theta_1) = \mathcal{R}_i(\theta_2)$). Suppose that learners update via full risk minimization. Then equilibria will have the form $(\alpha^{\text{eq}}, (0, 0))$ for any $\alpha^{\text{eq}} \in \Delta_2^2$.

Then starting from any (α^0, Θ^0) with a $\delta_\alpha, \delta_\theta$ ball of any equilibrium $(\alpha^{\text{eq}}, \Theta^{\text{eq}})$,

$$\alpha^1 = \nu(\alpha^0, \Theta^0), \quad \Theta^1 = (0, 0)$$

at which point the system is in a new equilibrium, since any allocation α is a fixed point when $\Theta = (0, 0)$ so $\alpha^t = \alpha^1$ and $\Theta^t = \Theta^1$ for all t . We have that $\|\Theta^{\text{eq}} - \Theta_0\| = 0$ and

$$\|\alpha^{\text{eq}} - \alpha^1\| = \|\nu(\alpha^{\text{eq}}, \Theta^{\text{eq}}) - \nu(\alpha^0, \Theta^0)\|$$

By the assumption of Lipschitzness, this distance will scale linearly in $\delta_\alpha, \delta_\theta$ so the definition of stability is satisfied for δ chosen proportionally to ϵ depending on the Lipschitz constant of ν .

In this example, any perturbation converges to a new fixed point within one time step. The continuity of the update functions ensures that the new fixed point is within a bounded distance of the original, satisfying the definition of stability. This example is not asymptotically stable: the allocation does not convergence back to the original point.

Example D.5 (Asymptotically stable segmentation). Consider the subpopulation and learner update rules as in the prior example, with one amendment. When $\mathcal{R}_i(\theta_1) = \mathcal{R}_i(\theta_2)$, subpopulation 1 re-allocates half of its mass from learner 2 to learner 1, and while subpopulation 2 re-allocates half its mass from learner 1 to learner 2. Thus the subpopulation update can be written as

$$\alpha_{1,:}^{t+1} = \begin{cases} \nu_1(\alpha_{1,:}^t, \Theta^t) & \mathcal{R}_1(\theta_1) \neq \mathcal{R}_1(\theta_2) \\ \begin{bmatrix} 1 & 1/2 \\ 0 & 1/2 \end{bmatrix} \alpha_{1,:}^t & \mathcal{R}_1(\theta_1) = \mathcal{R}_1(\theta_2) \end{cases}, \quad \alpha_{2,:}^{t+1} = \begin{cases} \nu_2(\alpha_{2,:}^t, \Theta^t) & \mathcal{R}_2(\theta_1) \neq \mathcal{R}_2(\theta_2) \\ \begin{bmatrix} 0 & 1/2 \\ 1 & 1/2 \end{bmatrix} \alpha_{2,:}^t & \mathcal{R}_2(\theta_1) = \mathcal{R}_2(\theta_2) \end{cases}$$

The only equilibrium has α^{eq} segmented with subpopulation i associated to learner i for $i = 1, 2$ and $\Theta^{\text{eq}} = (0, 0)$. It is straightforward to see that this is an asymptotically stable equilibrium, since for any $a \in \Delta_2$,

$$\lim_{t \rightarrow \infty} \begin{bmatrix} 1 & 1/2 \\ 0 & 1/2 \end{bmatrix}^t a = \begin{bmatrix} 1 & 1 \\ 0 & 0 \end{bmatrix} a = \begin{bmatrix} 1 \\ 0 \end{bmatrix} \quad \text{and} \quad \lim_{t \rightarrow \infty} \begin{bmatrix} 0 & 1/2 \\ 1 & 1/2 \end{bmatrix}^t a = \begin{bmatrix} 0 & 0 \\ 1 & 1 \end{bmatrix} a = \begin{bmatrix} 0 \\ 1 \end{bmatrix}.$$

Example D.6 (Asymptotically stable balanced market). Consider a setting similar to the previous example except that when $\mathcal{R}_i(\theta_1) = \mathcal{R}_i(\theta_2)$, subpopulation i moves half the mass from group 1 to group 2 and half the mass from group 2 to group 1

for all i . Then the subpopulation update can be written as

$$\alpha_{1,:}^{t+1} = \begin{cases} \nu_1(\alpha_{1,:}^t, \Theta^t) & \mathcal{R}_1(\theta_1) \neq \mathcal{R}_1(\theta_2) \\ \begin{bmatrix} 1/2 & 1/2 \end{bmatrix} \alpha_{1,:}^t & \mathcal{R}_1(\theta_1) = \mathcal{R}_1(\theta_2) \end{cases}, \quad \alpha_{2,:}^{t+1} = \begin{cases} \nu_2(\alpha_{2,:}^t, \Theta^t) & \mathcal{R}_2(\theta_1) \neq \mathcal{R}_2(\theta_2) \\ \begin{bmatrix} 1/2 & 1/2 \end{bmatrix} \alpha_{2,:}^{t+1} & \mathcal{R}_2(\theta_1) = \mathcal{R}_2(\theta_2) \end{cases}$$

The only equilibrium has $\alpha_{i,:}^{\text{eq}} = [1/2, 1/2]$ for $i = 1, 2$ and $\Theta^{\text{eq}} = (0, 0)$. It is straightforward to see that this is an asymptotically stable equilibrium, since for any $a \in \Delta_2$,

$$\lim_{t \rightarrow \infty} \begin{bmatrix} 1/2 & 1/2 \\ 1/2 & 1/2 \end{bmatrix}^t a = \begin{bmatrix} 1/2 & 1/2 \\ 1/2 & 1/2 \end{bmatrix} a = \begin{bmatrix} 1/2 \\ 1/2 \end{bmatrix}.$$

Example D.7 (Unstable balanced market). Suppose that subpopulation allocations follow a projected gradient descent update for all i :

$$\alpha_{i1}^{t+1} = \text{Proj}_{[0,1]}(\alpha_{i1}^t - \gamma(\mathcal{R}_i(\theta_1) - \mathcal{R}_i(\theta_2)))$$

and $\alpha_{i2} = 1 - \alpha_{i1}$. Further suppose learners update with gradient descent:

$$\theta_j^{t+1} = \theta_j^t - \frac{1}{2\sqrt{t}} \nabla \bar{\mathcal{R}}_j^{\text{learner}}(\alpha_{:,j}^t, \theta_j^t) = \sqrt{\frac{t}{t+1}} \theta_j^t$$

Both rules are risk minimizing in the limit (note that $\theta_j^t = \frac{1}{\sqrt{t}} \theta_j^0$) and have a continuum of equilibria at any $\alpha^{\text{eq}} \in \Delta_m^n$ and $\Theta^{\text{eq}} = (0, 0)$. However, we show that the equilibria are not stable. Consider the initial condition $(\alpha^{\text{eq}}, (\delta_\theta, 0))$. We have that

$$\alpha_{i1}^{t+1} = \text{Proj}_{[0,1]} \left(\alpha_{i1}^0 - \gamma \delta_\theta^2 \sum_{k=1}^t \frac{1}{k} \right) \rightarrow 0 \quad \text{as } t \rightarrow \infty.$$

No matter how small the perturbation δ_θ is, the summation increases with t and participation will converge all weight to learner 2. A similar argument shows that perturbations exist that will send all participation to learner 1.

In this example, the learners update slowly. Despite eventual convergence to the minimizing parameter, the accumulating error causes the participation allocation to shift completely to the unperturbed learner, precluding stability.

D.3 Social Welfare

We begin with a somewhat generic example with $m = 2$ and $n = 3$ that illustrates the difference between stable equilibria and social welfare optima.

Example D.8 (Stability vs. optimality). Consider three subpopulations $i \in \{1, 2, 3\}$ with risks $\|\theta - \phi_i\|_2^2$, sizes β_i , and two learners $j \in \{1, 2\}$. Suppose that the α^{eq} is such that the subpopulations are partitioned into $\{1\}$ and $\{2, 3\}$. Then we have that

$$\theta_1^{\text{eq}} = \phi_1, \quad \theta_2^{\text{eq}} = \frac{\beta_2}{\beta_2 + \beta_3} \phi_2 + \frac{\beta_3}{\beta_2 + \beta_3} \phi_3$$

By Theorem 4.6 this is stable if and only if

$$\|\phi_2 - \phi_3\|_2 \leq (\beta_2 + \beta_3) \min \left\{ \frac{\|\phi_2 - \phi_1\|_2}{\beta_3}, \frac{\|\phi_3 - \phi_1\|_2}{\beta_2} \right\}.$$

However, it is only social optimal if and only if ϕ_2 and ϕ_3 are relatively close to each other than to ϕ_1 , i.e.

$$\|\phi_2 - \phi_3\|_2 \leq \min \{\|\phi_2 - \phi_1\|_2, \|\phi_3 - \phi_1\|_2\}.$$

The set of subpopulation optima $\{\phi_1, \phi_2, \phi_3\}$ satisfying the optimality condition are a subset of those satisfying the stability condition. As the difference between β_2 and β_3 becomes more extreme, the number of settings satisfying the stability but not optimality condition increases.

We use this generic example to illustrate a scenario in which the total risk can be arbitrarily high at a stable equilibria.

Example D.9. Suppose there are two learners and three subpopulations with sizes $\beta_1 = \beta_2 = \beta$ and $\beta_3 = 1 - 2\beta$ for some $0 < \beta < 1/2$. Consider the following: $\mathcal{R}_1(\theta) = \theta^2$, $\mathcal{R}_2(\theta) = (\theta - 1)^2$, $\mathcal{R}_3(\theta) = (\theta - \frac{1-\beta}{1-2\beta} + \epsilon)^2$. The social welfare optimizing decision $\Theta^* = (1/2, \frac{1-\beta}{1-2\beta} - \epsilon)$ corresponds to total risk $\beta/2$. However, there is a stable equilibrium at $\Theta^{\text{eq}} = (0, 1 + \epsilon)$ with total risk $\beta + \frac{(\beta-\epsilon)^2}{1-2\beta}$. For $\beta \rightarrow 1/2$, the gap becomes arbitrarily large.

Finally, we present an example which illustrates that even in the single learner setting, the risk of a subpopulation can be arbitrarily worse than the total risk.

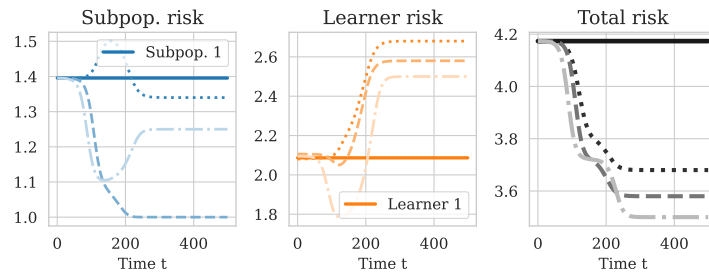
Example D.10 (Arbitrarily high risk for minority subpopulation). Consider two subpopulations with $\mathcal{R}_1(\theta) = \theta^2$ and $\mathcal{R}_2(\theta) = (\theta - \phi)^2$ with $\beta_1 = \beta$ and $\beta_2 = 1 - \beta$ and a single learner. The single equilibrium and total risk minimizer is $\theta_1 = (1 - \beta)\phi$ with total risk $\beta(1 - \beta)\phi^2$ and $\mathcal{R}_2(\theta^*) = \beta^2\phi^2$. The difference between the two quantities can be arbitrarily high as β gets close to 1.

E Additional Experiments with Noisy Dynamics

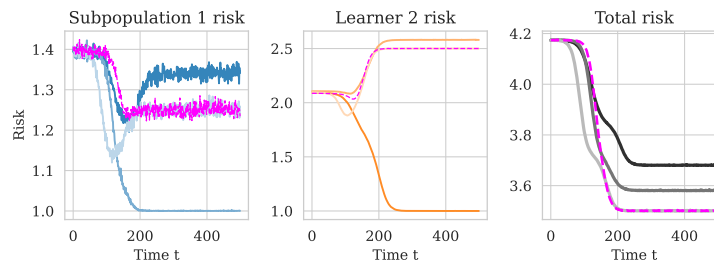
Figure 8a replicates Fig. 4a from the main text. The magenta-highlighted trajectory starts precisely at the unstable equilibrium, while the other three, initiated near this point, converge to the three possible split market equilibria, ordered by hue intensity: $\{(1,2), (3)\}$, $\{(2,3), (1)\}$, and $\{(1,3), (2)\}$. In Figure 8b, while sub-population dynamics remain as in (a), learner updates experience uncorrelated external perturbations, causing trajectories to be different from (a). Nevertheless, the long term dynamics gravitate near stable split equilibria. Figure 8c depicts learners updating decisions based on sampled empirical losses, with sub-populations adjusting participation based on aggregate empirical performance. The fact that each learner uses different samples from each sub-population adds sufficient un-correlated noise to create trajectories similar to when exogenous noise is added.

F Experimental Details

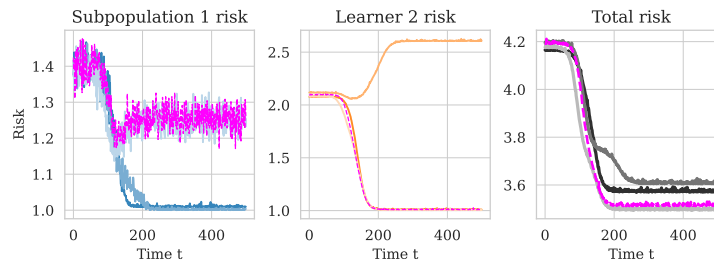
Full experimental details along with instructions for reproducing them can be found at <https://github.com/mcurmei627/MultiLearnerRiskReduction>. The experiments used Python 3.10 on a MacBook Pro 2019.



(a) Learner updates: **noiseless** one-step minimization of **population** loss. Sub-population updates: MWU w.r.t **population** loss



(b) Learner updates: **noisy** one-step minimization of **population** loss. Sub-population updates: MWU w.r.t **population** loss



(c) Learner updates: **noiseless** one-step minimization of **empirical** loss. Sub-population updates: MWU w.r.t **empirical** loss

Figure 8: Noisy dynamics