



Mathematics of Operations Research

Publication details, including instructions for authors and subscription information:
<http://pubsonline.informs.org>

A Mean Field Competition

Marcel Nutz, Yuchong Zhang

To cite this article:

Marcel Nutz, Yuchong Zhang (2019) A Mean Field Competition. Mathematics of Operations Research 44(4):1245-1263. <https://doi.org/10.1287/moor.2018.0966>

Full terms and conditions of use: <https://pubsonline.informs.org/Publications/Librarians-Portal/PubsOnLine-Terms-and-Conditions>

This article may be used only for the purposes of research, teaching, and/or private study. Commercial use or systematic downloading (by robots or other automatic processes) is prohibited without explicit Publisher approval, unless otherwise noted. For more information, contact permissions@informs.org.

The Publisher does not warrant or guarantee the article's accuracy, completeness, merchantability, fitness for a particular purpose, or non-infringement. Descriptions of, or references to, products or publications, or inclusion of an advertisement in this article, neither constitutes nor implies a guarantee, endorsement, or support of claims made of that product, publication, or service.

Copyright © 2019, INFORMS

Please scroll down for article—it is on subsequent pages



With 12,500 members from nearly 90 countries, INFORMS is the largest international association of operations research (O.R.) and analytics professionals and students. INFORMS provides unique networking and learning opportunities for individual professionals, and organizations of all types and sizes, to better understand and use O.R. and analytics tools and methods to transform strategic visions and achieve better outcomes.

For more information on INFORMS, its publications, membership, or meetings visit <http://www.informs.org>

A Mean Field Competition

Marcel Nutz,^a Yuchong Zhang^b

^aDepartments of Statistics and Mathematics, Columbia University, New York, New York 10027; ^bDepartment of Statistical Sciences, University of Toronto, Toronto, Ontario M5S 3G3, Canada

Contact: mnutz@columbia.edu,  <http://orcid.org/0000-0003-2936-2315> (MN); yuchong.zhang@utoronto.ca,  <http://orcid.org/0000-0002-7687-2783> (YZ)

Received: August 3, 2017

Revised: April 10, 2018

Accepted: August 12, 2018

Published Online in Articles in Advance:
June 27, 2019

MSC2000 Subject Classification:

Primary: 91A13; secondary: 91B40, 93E20

OR/MS Subject Classification:

Primary: games/group decisions: nonatomic;
research and development: innovation;
secondary: probability

<https://doi.org/10.1287/moor.2018.0966>

Copyright: © 2019 INFORMS

Abstract. We introduce a mean field game with rank-based reward: competing agents optimize their effort to achieve a goal, are ranked according to their completion time, and are paid a reward based on their relative rank. First, we propose a tractable Poissonian model in which we can describe the optimal effort for a given reward scheme. Second, we study the principal–agent problem of designing an optimal reward scheme. A surprising, explicit design is found to minimize the time until a given fraction of the population has reached the goal.

Funding: M. Nutz is supported by an Alfred P. Sloan Fellowship and the National Science Foundation (NSF) [Grants DMS-1512900 and DMS-1812661]. Y. Zhang is supported by the NSF [Grant DMS-1714607] and thanks the Department of Statistics at Columbia University, where much of the research was carried out, for its generous support.

Keywords: mean field game • rank-based reward • optimal contract • R&D competition

1. Introduction

In this paper, we introduce two game-theoretic problems. The first one is a mean field game with infinitely many players that compete to obtain a reward. The second one is a principal–agent problem where the principal interacts with these agents; namely, the principal aims to distribute a given reward budget to the different ranks such as to minimize the time until the agents complete their task.

Let us think of the agents as independent research teams trying to develop a result or product in the same field. Following the literature on dynamic research and development (R&D) detailed below, attaining a result will be modeled as a binary event. At any time t , each agent chooses a research effort λ for which a quadratic instantaneous cost $c\lambda^2$ is to be paid, where $c > 0$ is assumed to be constant for the purpose of this introduction. In a Poissonian fashion, the agents' probability of reaching the goal in a small time interval Δt is then given by $\lambda\Delta t + o(\Delta t)$; in the R&D literature, λ is sometimes interpreted as the accumulation of private knowledge, and the goal is to file a patent. The agents are ranked according to their completion times and paid a reward $R(r)$ for rank r , where the reward scheme R is a given decreasing function.¹ At any time, the agents observe the fraction $\rho(t)$ of players that have already completed the task and thus which portion of R is still available; more precisely, the agents use feedback controls $\lambda(\rho(t))$. This rank-based coupling of the agents' optimization problems is a nonstandard example of a mean field interaction. We shall show that given R , this game has a unique Nash equilibrium when agents optimize the expectation of reward minus cost. In fact, this setting turns out to be very tractable: Theorem 1 provides explicit formulas for the equilibrium optimal control λ^* and the agents' value function. The value function is independent of the cost c , which, in the body of this paper, is also allowed to depend on the state $\rho(t)$ to model that the cost may diminish as more results become available.

The second problem is built on top of the first: because we have seen that any reward scheme R leads to a unique equilibrium between the agents, we can study the problem of a manager or policymaker who would like to advance research. More precisely, we aim to minimize the time T_α until a given fraction α of the population has completed their task. The principal has a fixed reward budget $B = \int_0^1 R(r) dr$ but may choose the shape of the decreasing function R —that is, how much reward to allocate to each rank. Quite surprisingly, the principal's optimization problem has an explicit but nontrivial solution (Theorem 2):

$$R^*(r) = \frac{B}{C'} \left\{ \frac{1}{\sqrt{2-r}} + \frac{1}{2} \log \frac{(1 + \sqrt{2-\alpha})(1 - \sqrt{2-r})}{(1 - \sqrt{2-\alpha})(1 + \sqrt{2-r})} \right\} \mathbf{1}_{[0,\alpha]}(r),$$

where C' is a constant such that the budget constraint is saturated. As can be seen in Figure 1, this function has two main features. The first one is a discontinuity at $r = \alpha$: a substantial amount is awarded to the last few relevant agents, but it is optimal to pay zero reward to the ranks after α . Although it is clearly important to incentivize the last agents that will complete the α fraction, these agents are not too discouraged by the fact that they may miss the rewarding ranks. The second feature is the shape of R^* on $[0, \alpha]$. A priori, it may not even be obvious whether it is better to provide a strictly decreasing reward compared with paying the same amount to the first α ranks. It turns out that R^* is decreasing, even if not very much so, and moreover, it is concave. Thus, the difference in reward between two equidistant ranks increases later in the game, apparently to incentivize the remaining agents to choose a higher effort, as shown on the second panel of Figure 1.

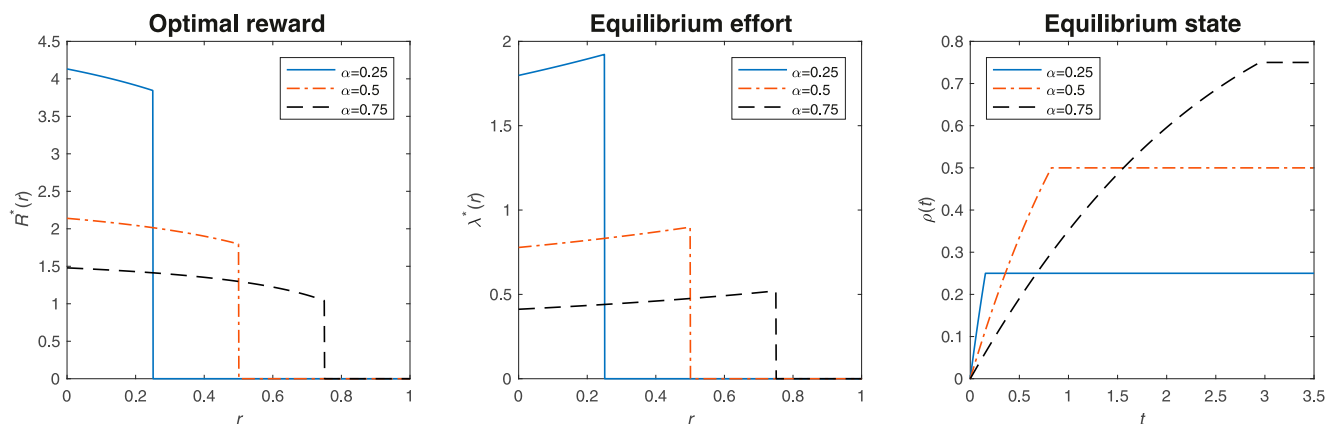
While formulating our game with a continuum of agents is convenient to get directly to our main results, a more fundamental justification for the mean field game is to study an N -player game for $N \rightarrow \infty$. Indeed, we establish that the N -player version of our two problems have unique solutions, albeit somewhat less explicit than in the mean field case. We show that the N -player equilibrium converges to the mean field limit; that is, the value functions and the optimal feedback controls converge (Theorem 4) if the given reward schemes converge. Moreover, the optimal reward schemes for the principal and the corresponding expected completion times converge (Theorem 5). The analysis for finite N also allows us to study size effects that cannot be observed in the mean field limit and thus are rarely addressed for large population games. In particular, we shall observe in Section 4.3 that with a fixed per capita budget, an increase in population size adversely affects the principal: the minimal expected completion time for a fixed target proportion α is increasing in N .

1.1. Literature

Dynamic competitions (also called *races*) are classic in the economics literature. An early reference related to our paper is Reinganum [34], which discusses an N -player dynamic game of R&D and patent protection. Rewards are paid at a fixed time horizon, and two cases are studied: either only the first-ranked player is rewarded (perfect patent protection) or the subsequent ranks receive a positive, smaller reward; however, only the case of an identical reward for all “imitators” is considered. Malueg and Tsutsui [29] extend Reinganum’s setting with a hazard rate that represents changes in the difficulty of the research project, an aspect we have incorporated differently by using a state-dependent cost that can model how the increase in public knowledge affects the project. Harris and Vickers [19] and Grossman and Shapiro [17] focus on the strategic interactions between agents in multiperiod two-player games. A recent work in this area is that of Cao [6], which studies a continuous-time, continuous-state version of a model in Harris and Vickers [19]. See also Cao [6], Malueg and Tsutsui [29], Taylor [40], Hörner [20], and Moscarini and Smith [30] for related contributions and further references in this literature.

Mean field games were introduced by Lasry and Lions [25, 26, 27] and Huang et al. [22, 23] to study Nash equilibria in the limiting regime where the number of players tends to infinity and interactions take place through the empirical distribution of the private states; we refer to Guéant et al. [18], Bensoussan et al. [4], and Carmona and Delarue [8, 9] for background on mean field games. Because the impact of an individual player on the aggregate distribution is negligible, finding a Nash equilibrium reduces to solving a stochastic optimization problem for a representative player against a fixed environment, together with a consistency condition. This can be justified rigorously by formulating a game with a continuum of players and by showing

Figure 1. (Color online) The optimal reward scheme R^* and the corresponding equilibrium effort λ^* and state process ρ for three different cutoff values α , with $B = 1$ and $c \equiv 1$.



that it is the limit of an N -player game. Convergence is often shown backward; that is, the mean field equilibrium is shown to provide an ε -Nash equilibrium for the N -player game. The forward convergence of the N -player equilibrium to the mean field equilibrium is typically more difficult to prove. For standard, diffusion-driven mean field games, this has been accomplished recently in the seminal work of Cardaliaguet et al. [7]. Although our game is of a different form, its tractability allows us to give an elementary yet nontrivial proof of the forward convergence; one feature in common with Cardaliaguet et al. [7] is that we use feedback controls. In a finite-state setting, forward convergence was shown by Gomes et al. [16] for a small time horizon and more recently by Bayraktar and Cohen [1] for an arbitrary finite time horizon.

Competitions (i.e., rank-based rewards) are a classic topic in contract theory, dating back to the work of Lazear and Rosen [28]. With applications from school grades to sports and business competitions, rank-order prize allocation is one of the most widely used relative performance evaluation criteria; we refer to Vojnović [41] for a detailed introduction and an extensive list of references. To the best of our knowledge, the only existing work on mean field games with rank-based reward is Bayraktar and Zhang [2], where players are ranked according to their terminal positions. The main aim of that work is to obtain abstract existence results for games with common noise via translation invariance, whereas in this work players are ranked according to exit times, and the focus is on specific properties of solutions (and, of course, the principal's problem). In a different way, exit times are used in the toy example of Guéant et al. [18]: "When does the meeting start?" Rank-based features are also studied in the literature on (uncontrolled) particle systems; one example is Shkolnikov [37]. Nadtochiy and Shkolnikov [31] consider particles interacting through hitting times. A different but related recent literature studies mean field games of timing where players directly choose stopping times; see Carmona et al. [11], Bertucci [5], and Nutz [32]. See also Seel and Strack [35, 36], and Feng and Hobson [14, 15] for a model where finitely many agents aim to stop in the highest state but do not observe the competitors.

Continuous-time principal–agent problems with multiple agents have been studied by Koo et al. [24] and Elie and Possamaï [12] and have been extended to the mean field setting by Elie et al. [13] and Bensoussan et al. [3]. Although these works have not considered rank-based rewards, a common feature is the Stackelberg equilibrium: the principal designs a reward scheme that the agents take as an external input to form a Nash equilibrium among themselves. By contrast, in mean field games with a major player, as in Huang [21] or Carmona and Wang [10], a Nash equilibrium is formed collectively by the major and minor players. To the best of our knowledge, the convergence of the N -player principal–agent problem to the mean field limit has only been established in a simple example where the equilibrium controls are independent of N ; see Elie et al. [13].

The remainder of this paper is structured as follows. In Section 2, we determine the unique Nash equilibrium of the mean field competition for a continuum of players with a given reward scheme. On the strength of this result, Section 3 solves the associated principal–agent problem where the principal designs the reward scheme. In Section 4, we study the corresponding N -player problems, and Section 5 establishes their convergence as $N \rightarrow \infty$. Proofs are gathered in Appendix A, whereas Appendix B provides background on the exact law of large numbers used in Section 2.

2. Mean Field Game

Let $(I, \mathcal{I}, \mu)$ be an atomless probability space; each $i \in I$ is thought of as an agent. Moreover, let $(\Omega, \mathcal{F}, P)$ be another probability space, to be used as the sample space. Let $(Z^i)_{i \in I}$ be a family of exponential(1)-distributed random variables on Ω , which is essentially pairwise independent; that is, for μ -almost all $i \in I$, Z^i is independent of Z^j for μ -almost all $j \in I$. We assume that this family is defined on an extension of the product $(I \times \Omega, \mathcal{I} \otimes \mathcal{F}, \mu \otimes P)$ for which the exact law of large numbers holds, as detailed in Appendix B. Given a locally Lebesgue-integrable function $\theta : \mathbb{R} \rightarrow [0, \infty)$, we define

$$\tau_\theta^i = \inf \left\{ t : \int_0^t \theta(s) ds = Z^i \right\}; \quad (1)$$

then $(\tau_\theta^i)_{i \in I}$ are essentially pairwise independent, and their distribution corresponds to the first jump time of an inhomogeneous Poisson process with intensity θ . Below, the function θ is of the form $\theta = \lambda \circ \rho$ where λ is a function chosen by the agent, and ρ is a given function, and we shall find it convenient to write τ_λ^i for τ_θ^i despite the abuse of notation. If $\tau_\lambda^i \leq t$, we shall say that agent i has "arrived" by time t .

We define the set Λ of *admissible (feedback²) controls* as the set of piecewise Lipschitz continuous³ functions $\lambda : [0, 1) \rightarrow \mathbb{R}_+$. The next lemma introduces the state process that emerges if all agents use the control $\lambda \in \Lambda$.

Lemma 1. Let $\lambda \in \Lambda$ be an admissible feedback control. There exists a unique continuous function $\rho : \mathbb{R}_+ \rightarrow [0, 1)$ satisfying

$$\rho(t) = \int_0^t \lambda(\rho(s))(1 - \rho(s)) ds, \quad t \geq 0. \quad (2)$$

If all agents use the feedback control λ , then $\rho(t) = \mu\{i : \tau_\lambda^i(\omega) \in [0, t]\}$ P -a.s. as well as $\rho(t) = P\{\tau_\lambda^i \in [0, t]\}$ μ -a.s.; that is, $\rho(t)$ is both the proportion of agents that have arrived by time t and the probability that any given agent i has arrived by time t .

Next, we fix a cost coefficient $c : [0, 1] \rightarrow (0, \infty)$, that is assumed to be Lipschitz continuous (thus, c and $1/c$ are bounded). Moreover, we fix a reward scheme $R : [0, 1] \rightarrow \mathbb{R}_+$ that is assumed to be decreasing, piecewise Lipschitz continuous, and left-continuous at $r = 1$. We interpret $R(r)$ as the reward paid to an agent arriving at rank r .

Let us now consider the control problem of a given agent i before arriving, assuming that all other agents use $\lambda_0 \in \Lambda$, which induces a state process $\rho = \rho_{\lambda_0}$ by (2). The objective of agent i is to choose $\lambda \in \Lambda$ such as to maximize

$$J(\lambda; \lambda_0) = E \left[R(\rho(\tau_\lambda)) - \int_0^{\tau_\lambda} c(\rho(t)) \lambda(\rho(t))^2 dt \right],$$

where $\tau_\lambda = \tau_\lambda^i$ is the arrival time of the given agent for the control λ . Thus, the agent is rewarded according to the rank but pays a quadratic cost of effort with coefficient c . We use the convention $\rho(\infty) := 1$, meaning that agents who never arrive are paid the reward $R(1)$. If $\lambda \in \Lambda$ attains $\sup_{\lambda \in \Lambda} J(\lambda; \lambda_0)$, we say that λ is an *optimal control given λ_0* . Moreover, we introduce the value function

$$v(r) = v(r; \lambda_0) = \sup_{\lambda \in \Lambda} E \left[R(\rho(\tau_\lambda)) - \int_0^{\tau_\lambda} c(\rho(t)) \lambda(\rho(t))^2 dt \middle| \rho(0) = r \right], \quad (3)$$

where $d\rho(t) = \lambda_0(\rho(t))(1 - \rho(t))dt$. Note that $v(0) = \sup_{\lambda \in \Lambda} J(\lambda; \lambda_0)$.

If λ is an optimal control given λ_0 , we call λ an *equilibrium optimal control* and the induced ρ the corresponding *equilibrium state process*. This is a Nash equilibrium: if all other players use the feedback control λ , the state evolves according to ρ by Lemma 1 (recall that μ is atomless), and then λ is an optimal control for our fixed player.

Theorem 1. Let R be a reward scheme. Then there exists a unique (almost everywhere [a.e.]) equilibrium optimal control $\lambda^* \in \Lambda$, given by

$$\lambda^*(r) = \frac{R(r) - \frac{1}{2\sqrt{1-r}} \int_r^1 \frac{R(y)}{\sqrt{1-y}} dy}{2c(r)}, \quad r \in [0, 1), \quad (4)$$

and the corresponding equilibrium state process ρ is determined by (2) with $\lambda = \lambda^*$. In equilibrium, the value function of any agent before arriving is

$$v(r) = \frac{1}{2\sqrt{1-r}} \int_r^1 \frac{R(y)}{\sqrt{1-y}} dy, \quad r \in [0, 1). \quad (5)$$

Let us remark that although the piecewise Lipschitz requirement is mainly for convenience, the continuity of R at $r = 1$ is more essential in providing existence. The following example exhibits a phenomenon that is familiar in infinite-horizon optimal stopping problems.

Example 1. Suppose that $c \equiv 1$ and $R = \mathbf{1}_{[0,1]}$; that is, the reward is one for agents arriving in finite time and zero for those who never arrive. Then, using the constant control $\lambda \equiv \varepsilon > 0$ yields an exponential arrival time with $E[\tau] = 1/\varepsilon$, and thus the expected reward is

$$E \left[R(\rho(\tau)) - \int_0^\tau \varepsilon^2 dt \right] = 1 - \varepsilon.$$

As a result, the value function satisfies $v(r) = 1 = R(r)$ for all $r < 1$. But because $\lambda \equiv 0$ yields zero reward and any other control has positive cost, this value is not attained: there is no optimal control and thus no equilibrium in the above sense.

Remark 1. We see from (5) that the equilibrium value function is independent of the cost coefficient c . This can also be understood directly by expressing (3) as an integral over the ranks, using the Kolmogorov equation (2) and the change-of-variable formula:

$$v(r_0) = \sup_{\lambda \in \Lambda} E \left[R(\rho(\tau_\lambda)) - \int_{\rho(0)}^{\rho(\tau_\lambda)} c(r) \lambda(r) \frac{dr}{1-r} \middle| \rho(0) = r_0 \right].$$

Indeed, $\rho(\tau_\lambda)$ is independent of λ —when all agents use the same control, their ranking is given by the ranking of the Z^i . By contrast, $c\lambda \in \Lambda$ if and only if $\lambda \in \Lambda$, and hence v is independent of c in equilibrium. Intuitively, a higher cost leads to a smaller optimal effort, but because this holds for all agents, the equilibrium state ρ is slowed down to the extent that the reduced effort results in the same reward.

The equilibrium value function of any agent has a surprising interpretation: it can be compared with a deal where the agent pays no cost for his (constant) effort but is given the handicap of running at half the intensity of the competitors.

Proposition 1. *The equilibrium value function v of (5) coincides with the value function of an agent whose effort is fixed at $\lambda \equiv \lambda_0 \in (0, \infty)$ and is charged zero cost, whereas all other agents use $\lambda \equiv 2\lambda_0$:*

$$v(r) = E[R(\rho(\tau)) | \rho(0) = r], \quad \text{where } \tau \sim \text{Exp}(\lambda_0) \text{ and } \rho'(t) = 2\lambda_0(1 - \rho(t)).$$

In particular, $v(0) = E[R(1 - e^{-2\tau})]$ for $\tau \sim \text{Exp}(1)$.

The following result shows that the unique equilibrium of Theorem 1 is stable with respect to the reward scheme.

Proposition 2. *Let R_n, R be reward schemes such that $R_n \rightarrow R$ pointwise. Then the corresponding equilibrium optimal controls also converge pointwise, whereas the equilibrium value functions and state processes converge uniformly.*

2.1. Examples with Closed-Form Solutions

In this section, we present a family of explicitly solvable examples. Given a total reward budget $B = \int_0^1 R(r) dr \geq 0$, the family has a cutoff parameter $\alpha \in (0, 1]$ indicating that no reward will be paid to agents ranked lower than α , as well as a shape parameter $q \geq 0$. The general form is then given by

$$R(r) = \kappa(1-r)^q \mathbf{1}_{[0, \alpha]}(r), \quad \kappa = \frac{B(1+q)}{1 - (1-\alpha)^{1+q}};$$

the constant κ is chosen such that $B = \int_0^1 R(r) dr$. We notice that a larger value of q indicates that a larger portion of the reward budget is paid to highly ranked player, whereas $q = 0$ corresponds to a uniform distribution of the reward among the top α ranks. For such a reward, the value function and the optimal effort of Theorem 1 admit closed-form solutions:

$$v(r) = \frac{\kappa}{(1+2q)} \left((1-r)^q - (1-\alpha)^q \sqrt{\frac{1-\alpha}{1-r}} \right)^+,$$

$$\lambda^*(r) = \mathbf{1}_{\{r \leq \alpha\}} \frac{\kappa}{2c(r)(1+2q)} \left(2q(1-r)^q + (1-\alpha)^q \sqrt{\frac{1-\alpha}{1-r}} \right).$$

In the boundary case $\alpha = 1$, the Lipschitz assumption of Theorem 1 is not satisfied when $0 < q < 1$. However, one can check the indicated formulas by direct computation in this case.

In general, the cumulative distribution function $F_{\tau_{\lambda^*}}(t) = \rho(t)$ of any agent's equilibrium completion time can be computed numerically by solving the Kolmogorov equation (2) for ρ . Inverting the equilibrium state process also gives rise to the quantile $T_\beta = \inf\{t : \rho(t) \geq \beta\}$ —that is, the time until a β -proportion of the players has reached the goal. In the following special cases, these quantities can be obtained in closed form.

2.1.1. Power Reward Without Cutoff. This case corresponds to $\alpha = 1$, where we also assume that the cost c is constant. Then the above formulas specialize to

$$v(r) = \frac{B(1+q)}{1+2q} (1-r)^q, \quad \lambda^*(r) = \frac{Bq(1+q)}{c(1+2q)} (1-r)^q,$$

and we can also solve for

$$F_{\tau_{\lambda^*}}(t) = \rho(t) = 1 - \left(1 + \frac{Bq^2(1+q)}{c(1+2q)}t\right)^{-\frac{1}{q}}, \quad T_\beta = \frac{c(1+2q)}{Bq^2(1+q)}[(1-\beta)^{-q} - 1].$$

We see that the equilibrium value v is decreasing in q . That is, each individual is worse off if the reward scheme heavily favors the highly ranked players; this can be attributed to the cost caused by the large effort level λ^* in the beginning of the competition. We also observe that λ^* is decreasing in r so that agents decrease their effort once the higher ranks are filled.

2.1.2. Uniform Reward with Cutoff. This case corresponds to $q = 0$, and we again assume that the cost c is constant. The general formulas now specialize to

$$v(r) = \frac{B}{\alpha} \left(1 - \sqrt{\frac{1-\alpha}{1-r}}\right)^+, \quad \lambda^*(r) = 1_{\{r \leq \alpha\}} \frac{B}{2c\alpha} \sqrt{\frac{1-\alpha}{1-r}}.$$

We also have $F_{\tau_{\lambda^*}}(t) = \rho(t) = 1 - \left(1 - \frac{B\sqrt{1-\alpha}}{4c\alpha}t\right)^2$ for $t \leq T_\alpha$ and $F_{\tau_{\lambda^*}}(t) = \rho(t) = \alpha$ for $t > T_\alpha$, where

$$T_\alpha = \frac{4c\alpha(1 - \sqrt{1-\alpha})}{B\sqrt{1-\alpha}}, \quad (6)$$

and then the general quantile is $T_\beta = 4c\alpha(1 - \sqrt{1-\beta})/(B\sqrt{1-\alpha})$ for $\beta \leq \alpha$ and $T_\beta = \infty$ for $\beta > \alpha$. In contrast to the case $\alpha = 1$, we see that λ^* is increasing in r for $r \leq \alpha$: as the race progresses, the agents compete for the remaining reward and increase their effort up to the time when an α -proportion of agents has reached the goal, and then the remaining players give up (see Figure 2).

2.2. Staircase Reward

Consider a reward scheme R and a cost coefficient c of the staircase form

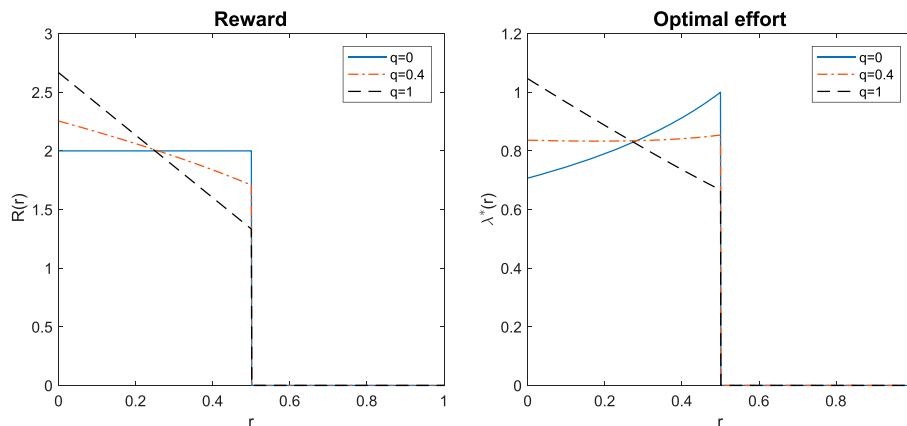
$$R = R_1 \mathbf{1}_{[r_0, r_1]} + \sum_{j=2}^n R_j \mathbf{1}_{(r_{j-1}, r_j]}, \quad c = c_1 \mathbf{1}_{[r_0, r_1]} + \sum_{i=2}^n c_i \mathbf{1}_{(r_{j-1}, r_j]},$$

where $R_1 \geq R_2 \geq \dots \geq R_n \geq 0$ and $0 = r_0 < r_1 < \dots < r_n = 1$ are constants. Equations (4) and (5) then yield that for $r_{j-1} < r \leq r_j$,

$$v(r) = R_j + \frac{1}{\sqrt{1-r}} \left[-R_j \sqrt{1-r_j} + \sum_{k=j+1}^n R_k (\sqrt{1-r_{k-1}} - \sqrt{1-r_k}) \right],$$

$$\lambda^*(r) = \frac{1}{2c_j \sqrt{1-r}} \left[R_j \sqrt{1-r_j} - \sum_{k=j+1}^n R_k (\sqrt{1-r_{k-1}} - \sqrt{1-r_i}) \right].$$

Figure 2. (Color online) Optimal effort under power reward with cutoff $\alpha = 0.5$, assuming $B = 1$ and $c \equiv 1$.



We claim that the equilibrium state ρ is given by

$$\rho(t) = 1 - \left(\sqrt{1 - r_{j-1}} - \frac{A_j}{4c_j}(t - t_{j-1}) \right)^2, \quad t_{j-1} \leq t \leq t_j, \quad (7)$$

where $A_j = R_j\sqrt{1 - r_j} - \sum_{k=j+1}^n R_k(\sqrt{1 - r_{k-1}} - \sqrt{1 - r_k})$, and t_j is recursively defined by $t_j = t_{j-1} + \frac{4c_j}{A_j}(\sqrt{1 - r_{j-1}} - \sqrt{1 - r_j})$ and $t_0 = 0$. This is to be read with the convention that $1/0 = \infty$; indeed, we have $A_j = 0$ and $t_j = \infty$ if (and only if) $R_j = R_{j+1} = \dots = R_n$. To see (7), we may solve the ordinary differential equation (ODE) (2) successively for each interval $[r_{j-1}, r_j]$. Let $t_0 = 0$. Suppose that we have already found $t_0, \dots, t_{j-1}$ and $\rho(t)$ for $t \in [0, t_{j-1}]$. Then the ODE on the j th interval reads as $\rho'(t) = \frac{A_j}{2c_j}\sqrt{1 - \rho(t)}$ with initial condition $\rho(t_{j-1}) = r_{j-1}$, and the solution is given by (7), whereas t_j is determined through the condition $\rho(t_j) = r_j$.

Finally, let $\beta \in (0, 1]$ be given. By adding β to the grid if necessary, we may assume without loss of generality that $\beta = r_{j_0}$ for some $j_0 \in \{1, \dots, n\}$, and then the β -quantile is $T_\beta = t_{j_0} = \sum_{j=1}^{j_0} \frac{4c_j}{A_j}(\sqrt{1 - r_{j-1}} - \sqrt{1 - r_j})$.

3. Mean Field Principal–Agent Problem

We have seen that for a given reward scheme R , there exists a unique (deterministic) equilibrium state ρ , and thus for $\alpha \in (0, 1]$, the time

$$T_\alpha(R) = \inf\{t \geq 0 : \rho(t) \geq \alpha\} \in (0, \infty]$$

is deterministic and well defined. This is the time until an α -proportion of the population has reached the goal, or equivalently, the α -quantile of the distribution of the equilibrium arrival time τ^* .

In this section, we fix $\alpha \in (0, 1)$ and the total reward budget $B > 0$, and we ask for a reward scheme R that minimizes $T_\alpha(R)$ subject to the constraint that $\int_0^1 R(r) dr \leq B$. This corresponds to a principal–agent problem in the second-best sense: the planner can set the reward for the agents but cannot dictate their choice of controls. The principal considers her project completed when an α -proportion of the agents have reached their goal, and she aims to find the minimal completion time

$$T_\alpha^* = \inf_{R \in \mathcal{R}: \int R(r) dr \leq B} T_\alpha(R), \quad (8)$$

where $\mathcal{R}$ is the set of all reward schemes. We remark that for $\alpha = 1$, we have $T_\alpha(R) = \infty$ for all R , whence we do not consider this case. By contrast, $T_\alpha^* < \infty$ for all $\alpha \in (0, 1)$, because this is already accomplished by the uniform reward R with cutoff at α (see (6)).

An additional assumption on the cost coefficient c is needed for our result:

$$r \mapsto \frac{c(r)(1 - r)}{2 - r} \quad \text{is decreasing.} \quad (9)$$

This assumption is discussed in more detail in Remark 2. The solution to the principal's problem is then given as follows.

Theorem 2. Let c satisfy (9). Given a reward budget $B > 0$ and $\alpha \in (0, 1)$, there is an a.e. unique optimal reward scheme R^* attaining the minimal completion time T_α^* of (8), given by

$$R^*(r) = \frac{B}{C} \left\{ \sqrt{\frac{c(r)}{2 - r}} + \frac{1}{2} \int_r^\alpha \frac{1}{1 - s} \sqrt{\frac{c(s)}{2 - s}} ds \right\} \mathbf{1}_{[0, \alpha]}(r), \quad (10)$$

and the minimal completion time is

$$T_\alpha^* = \frac{4C^2}{B}, \quad \text{where} \quad C = \frac{1}{2} \int_0^\alpha \frac{\sqrt{c(r)(2 - r)}}{1 - r} dr. \quad (11)$$

The corresponding equilibrium effort is

$$\lambda^*(r) = \frac{B}{2C} \frac{1}{\sqrt{(2 - r)c(r)}} \mathbf{1}_{[0, \alpha]}(r).$$

In the particular case where the cost c is constant, we have

$$\begin{aligned} R^*(r) &= \frac{B}{C'} \left\{ \frac{1}{\sqrt{2-r}} + \frac{1}{2} \log \frac{(1 + \sqrt{2-\alpha})(1 - \sqrt{2-r})}{(1 - \sqrt{2-\alpha})(1 + \sqrt{2-r})} \right\} \mathbf{1}_{[0,\alpha]}(r), \\ T_\alpha^* &= \frac{4cC'^2}{B}, \\ C' &= \frac{C}{\sqrt{c}} = \sqrt{2} - \sqrt{2-\alpha} + \frac{1}{2} \log \frac{(1 + \sqrt{2-\alpha})(1 - \sqrt{2})}{(1 - \sqrt{2-\alpha})(1 + \sqrt{2})}. \end{aligned}$$

Figure 1 shows R^* , λ^* , and ρ for constant cost coefficient c . As discussed, the key features of R^* are the strict decrease and concavity on $[0, \alpha]$ and the discontinuity at α . The equilibrium effort λ^* is strictly increasing on $[0, \alpha]$. For general c , the product $\sqrt{c}\lambda^*$ is increasing, but λ^* need not be.

Remark 2. Assumption (9) is satisfied in particular if c is decreasing, which certainly holds in the applications we have in mind. If we suppose for simplicity that c is differentiable, the assumption is equivalent to the derivative c' satisfying $c'(r) \leq \frac{c(r)}{(2-r)(1-r)}$, for which a sufficient condition is that $c'(r) \leq c(r)/2$. Thus, an increase of c is permissible if it is limited relative to the level of c . When the assumption is not satisfied, (10) no longer describes the solution to the principal–agent problem; in fact, (10) is not a decreasing function and hence not a reward scheme. The proof of Theorem 2 shows that finding an optimal reward scheme can still be phrased as a convex optimization problem; however, the monotonicity constraint on R is now binding, which is an obstruction to finding an explicit solution.

We may also ask the reverse question: given $\alpha \in (0, 1)$ and a desired completion time $T > 0$, what is the minimal budget enabling the principal to achieve T ? The answer follows from Theorem 2 by inverting (11).

Corollary 1. Let c satisfy (9). Given $\alpha \in (0, 1)$ and $T > 0$, the minimal budget enabling the principal to achieve a completion time $T_\alpha^* \leq T$ is $B^* = 4C^2/T$, where C is given by (11).

4. N-Player Problems

In this section, we study a version of the competition with finitely many players as well as a corresponding version of the principal–agent problem. The connections between these and the mean field formulations will be established in Section 5.

4.1. The N-Player Game

We consider a game with N players, where $N \geq 1$ is a fixed integer. At any time t , each player i observes the number n of players that have already arrived and chooses an effort level $\lambda_i(n) \in \mathbb{R}_+$. Suppose that player i uses the feedback control λ_i , and all other players use a feedback control λ_{-i} . We denote by $\xi_{\lambda_i, \lambda_{-i}}(t)$ the number of players that have arrived by time t ; that is,

$$\xi_{\lambda_i, \lambda_{-i}}(t) = \sum_{i=1}^N \mathbf{1}_{\{\tau^i \leq t\}},$$

where $\tau^i = \inf\{t \geq 0 : \int_0^t \lambda_i(\xi_{\lambda_i, \lambda_{-i}}(s)) ds = Z^i\}$ is the arrival time of player i , and $\tau^j = \inf\{t \geq 0 : \int_0^t \lambda_{-i}(\xi_{\lambda_i, \lambda_{-i}}(s)) ds = Z^j\}$ is the arrival time of player $j \neq i$, for some independent exponential random variables $\{Z^1, \dots, Z^N\}$ with unit rate. The existence of the state process is clear in this case; we may see $(\mathbf{1}_{\{\tau^i \leq t\}}, \xi_{\lambda_i, \lambda_{-i}})$ as a Markov pure jump process with values in $\{0, 1\} \times \{0, 1, \dots, N\}$. We emphasize that we now use the number rather than the fraction of arrived players as the state variable.

Let $R_n \in \mathbb{R}_+$ be the reward for finishing at the n th place; as before, we assume that $(R_n)_{1 \leq n \leq N}$ is decreasing, and we convene that R_N is paid to players that never arrive. Moreover, let $c_n > 0$ be the cost coefficient when n players have arrived. Then the objective of player i is to maximize

$$J_i(\lambda_i; \lambda_{-i}) = E \left[R_{\xi_{\lambda_i, \lambda_{-i}}(\tau^i)} - \int_0^{\tau^i} c_{\xi_{\lambda_i, \lambda_{-i}}(s)} \lambda_i^2(\xi_{\lambda_i, \lambda_{-i}}(s)) ds \right],$$

and λ is a (symmetric) equilibrium optimal control if $\arg \max_{\lambda_i} J_i(\lambda_i; \lambda) = \lambda$ for all i . For $0 \leq n \leq N-1$, the value function of player i before arriving is

$$v_n := \sup_{\lambda_i} E \left[R_{\xi(\tau^i)} - \int_0^{\tau^i} c_{\xi(s)} \lambda_i^2(\xi(s)) ds \mid \xi(0) = n \right],$$

where $\xi := \xi_{\lambda_i, \lambda_{-i}}$. We also convene that $v_N := R_N$.

Remark 3. The last player never arrives. Indeed, once $N-1$ players have arrived, the remaining player achieves the optimal value v_{N-1} by using the control $\lambda_i \equiv 0$, and in fact, $v_{N-1} = R_N = v_N$. This is due to the convention that R_N is paid to players that never arrive. On the other hand, this convention is necessary for the existence of an equilibrium because the same value is asymptotically achieved by using the control $\lambda_i \equiv \varepsilon$ with small $\varepsilon > 0$.

Proposition 3. The N -player game has a unique Nash equilibrium. The equilibrium value function $(v_n)_{0 \leq n \leq N}$ is the unique solution of the backward recursion

$$v_n = \frac{R_{n+1} + 2(N-n-1)v_{n+1}}{1 + 2(N-n-1)}, \quad 0 \leq n \leq N-1; \quad v_N = R_N. \quad (12)$$

The unique equilibrium optimal control is

$$\lambda^*(n) = \frac{R_{n+1} - v_n}{2c_n}, \quad 0 \leq n \leq N-1. \quad (13)$$

4.2. The N -Player Principal-Agent Problem

Next, we consider the N -player version of the mean field principal-agent problem introduced in Section 3. Given $n_0 \in \{1, \dots, N-1\}$ and a (nonnegative, decreasing) reward scheme (R_n) , let

$$T_{n_0} = \inf\{t \geq 0 : \xi_{\lambda^*}(t) = n_0\}$$

be the (random) time until n_0 players have arrived, where the players use the unique equilibrium optimal control λ^* for (R_n) and the (fixed, positive) cost coefficients (c_n) (see (13)). Below we shall find it useful to write λ_n^* rather than $\lambda^*(n)$ whenever we are in the N -player setting.

Given the per capita⁴ reward budget $B > 0$, the principal chooses a reward scheme (R_n) such as to minimize the expected completion time ET_{n_0} subject to the budget constraint $\sum_{n=1}^N R_n \leq NB$. In analogy to (9), we shall assume that c^N satisfies

$$c_n^N \leq c_{n-1}^N \frac{(2N-2n+1)^2(2N-n-1)}{4(N-n-1)(N-n+1)(2N-n)}, \quad n < n_0; \quad (14)$$

again, this is satisfied, for example, when $n \mapsto c_n^N$ is constant or decreasing.

Theorem 3. Let c^N satisfy (14) and define $y_{n_0} = 0$:

$$y_n = \sqrt{\frac{c_n N(N-n-1)}{(N-n)(2N-n-1)}}, \quad n < n_0.$$

Given a per capita reward budget $B > 0$, there is a unique optimal reward scheme (R_n^*) attaining the minimal expected completion time $ET_{n_0}^*$ given by

$$R_n^* = \frac{B}{C} \left\{ y_{n-1} + \frac{1}{2} \sum_{k=n-1}^{n_0-1} \frac{y_k}{N-k-1} \right\} \mathbf{1}_{\{n \leq n_0\}},$$

and the minimal expected completion time is

$$ET_{n_0}^* = \frac{4C^2}{B}, \quad \text{where} \quad C = \frac{1}{2\sqrt{N}} \sum_{n=0}^{n_0-1} \sqrt{\frac{c_n(2N-n-1)}{(N-n)(N-n-1)}}.$$

The corresponding equilibrium optimal control is

$$\lambda_n^* = \frac{B}{2C} \sqrt{\frac{N(N-n-1)}{c_n(N-n)(2N-n-1)}} \mathbf{1}_{\{n < n_0\}}.$$

4.3. Size Effects

We conclude this section with a brief discussion of the influence of the population size N on the principal's problem in two different ways. To make the problems comparable, we assume that $c_n^N \equiv c$ is a constant independent of N .

1. First, we consider as above a principal with a given per capita budget B aiming to minimize the expected time until an α proportion of the population has arrived. The left panel in Figure 3 shows a negative size effect: the minimum expected completion time $ET_{n_0}^N$ for $n_0 = \lceil \alpha N \rceil$ is increasing in N ; that is, an increase in population size *adversely* affects the principal.

2. Second, we fix the total budget $K = NB$ and a completion head count n_0 . Then

$$ET_{n_0}^N = \frac{1}{NK} \left\{ \sum_{n=0}^{n_0-1} \sqrt{\frac{c}{1-n/N} \left(1 + \frac{1}{1-(n+1)/N} \right)} \right\}^2$$

is strictly decreasing in N with a limit equal to zero (right panel in Figure 3). That is, the principal who is aiming for a fixed number of completions benefits from an increase in the population size.

5. Convergence to the Mean Field

In this section, we show that the N -player competition and principal-agent problem converge to their mean field counterparts as $N \rightarrow \infty$.

5.1. Convergence of the N -Player Equilibrium

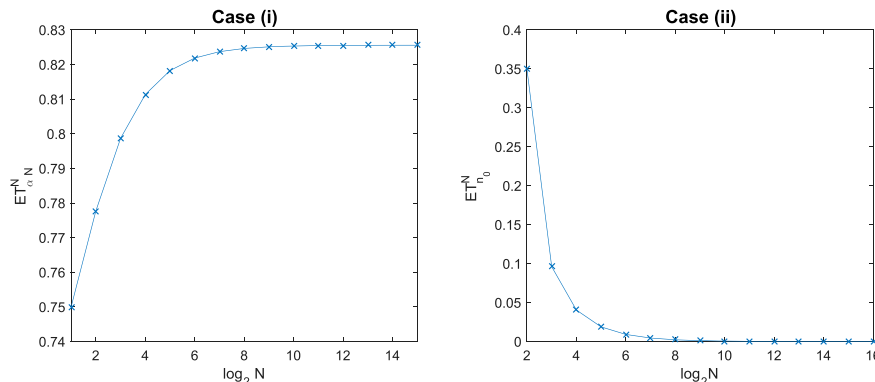
We consider the mean field setting of Section 2 with a fixed reward scheme R and cost coefficient c , as well as an N -player game with reward (R_n^N) and cost (c_n^N) . Our aim is to show that if $R^N \rightarrow R$ and $c^N \rightarrow c$ in a suitable sense, then the corresponding equilibria converge.

If we start with a reward scheme $R : [0, 1] \rightarrow \mathbb{R}_+$ in the mean field setting, an obvious choice for R^N is the sampling $R_n^N = R(\frac{n}{N})$. Because R is decreasing, we have $\frac{1}{N} \sum_{n=1}^N R_n^N \leq \int_0^1 R(r) dr$; that is, this discretization may (and typically will) reduce the cumulative reward. Another choice is the moving average $R_n^N = N \int_{(n-1)/N}^{n/N} R(y) dy$ which preserves the reward. Next, we introduce a condition designed to cover either of these choices and more.

Recall that $R : [0, 1] \rightarrow \mathbb{R}_+$ is decreasing, piecewise Lipschitz continuous, and left continuous at $r = 1$. Let $0 = r_0 < r_1 < \dots < r_m < r_{m+1} = 1$ be a finite partition of $[0, 1]$ such that R is Lipschitz on each interval (r_{i-1}, r_i) . For our results, we shall assume that

$$\sup_{r \in \bigcup_{i=1}^{m+1} I_i^N} |R_{\lceil rN \rceil}^N - R(r)| \leq \frac{K}{N} \quad (15)$$

Figure 3. (Color online) The left panel shows a negative size effect when fixing $\alpha = 0.5$ and $B = 1$. The right panel shows a positive size effect when fixing $n_0 = 3$ and $K = 2^5$. In both panels, $c \equiv 1$, and the minimal expected completion time is plotted against $\log_2 N$.



for some constant K independent of N , where $I_i^N = [r_{i-1} + 1/N, r_i - 1/N]$ for $i = 1, \dots, m$ and $I_{m+1}^N = [r_m + 1/N, 1]$. Similarly, we assume that⁵

$$\sup_{r \in [0,1]} |c_{[rN]}^N - c(r)| \leq \frac{K}{N}. \quad (16)$$

The next result establishes that the N -player equilibrium converges to the mean field equilibrium as $N \rightarrow \infty$. Thus, it gives a second justification to the mean field formulation, apart from the direct derivation with a continuum of players as in Section 2.

Theorem 4. Let R^N, R and c^N, c satisfy (15) and (16), and let v^N, v and λ^N, λ^* be the corresponding value functions and equilibrium optimal controls, respectively. Then

$$\sup_{r \in [0,1]} |v_{[rN]}^N - v(r)| = O(1/\sqrt{N}), \quad \sup_{r \in \bigcup_{i=1}^{m+1} I_i^N} |\lambda_{[rN]}^N - \lambda^*(r)| = O(1/\sqrt{N}).$$

Remark 4. If $R(r) = 0$ on $(\alpha, 1]$ for some $\alpha < 1$, then the ODE (A.2) is nondegenerate up to the boundary, and the convergence rates in Theorem 4 can be improved to $O(1/N)$ (see the proof of the theorem). In particular, this holds for the optimal reward scheme R^* of the principal's problem in Theorem 2.

5.2. Convergence and ε -Optimality for the Principal

Following Section 3, we fix a cost coefficient c satisfying (9), a target proportion $\alpha \in (0, 1)$, and a budget $B > 0$. We have seen in Theorem 2 that in the mean field setting, there exists a unique reward scheme R^* that attains the minimal (deterministic) time T_α^* until an α -proportion of the population has arrived. For the N -player situation with a given cost c^N satisfying (14) and per capita budget B , we have seen in Theorem 3 that there exists a unique reward scheme R^N minimizing the expected time $ET_{n_0}^N$ until n_0 players have arrived. In the following result, we consider $n_0 = \lceil \alpha N \rceil$ so that the proportion n_0/N tends to α and show that if $c^N \rightarrow c$, the expected completion time and the corresponding reward schemes converge as $N \rightarrow \infty$.

Theorem 5. Let $c^N \rightarrow c$ in the sense of (16). Then

$$|ET_{\lceil \alpha N \rceil}^N - T_\alpha^*| = O(1/N), \quad \sup_{r \in [0, \alpha]} |R_{[rN]}^N - R^*(r)| = O(1/N).$$

The next result addresses the convergence of the principal's problem in a different sense: it shows that if the principal applies (a discretization of) the optimal reward scheme R^* from the mean field setting in an N -player game with large N , rather than the precise optimal scheme R^N of Theorem 3, then R^* is still ε -optimal for minimization of the expected completion time.

Corollary 2. Let $c^N \rightarrow c$ in the sense of (16), and let $R^{(N)}$ be a discretization of R^* satisfying (15). Then the completion time $T_{\lceil \alpha N \rceil}^{(N)}$ of the N -player game under $R^{(N)}$ satisfies

$$|ET_{\lceil \alpha N \rceil}^{(N)} - ET_{\lceil \alpha N \rceil}^N| = O(1/N);$$

that is, the reward scheme $R^{(N)}$ is $O(1/N)$ -optimal for the N -player principal-agent problem.

Appendix A. Proofs

A.1. Proofs for Section 2

Proof of Lemma 1. The piecewise Lipschitz property of λ implies that (2) has a unique continuous solution ρ . This function is nonnegative, increasing, globally Lipschitz continuous, and continuously differentiable except at finitely many points (corresponding to the jumps to λ) where the left and right derivatives of ρ may disagree. Let $\bar{\rho}(t) = 1 - \exp(-\int_0^t \lambda(\rho(s)) ds)$. By the exact law of large numbers (Proposition 4) and (1),

$$\mu\{i : \tau_\lambda^i(\omega) \in [0, t]\} = P\{\omega : \tau_\lambda^i(\omega) \in [0, t]\} = \bar{\rho}(t)$$

holds almost surely. On the other hand, the derivative of $\bar{\rho}$ is seen to satisfy $\bar{\rho}'(t) = \lambda(\rho(t))(1 - \bar{\rho}(t))$ a.e., so the uniqueness of the exponential ODE yields $\bar{\rho} = \rho$. $\square$

Proof of Theorem 1. (1) Suppose that $\bar{\lambda} \in \Lambda$ is an equilibrium optimal control; let ρ be the corresponding equilibrium state process, and let $\bar{v}$ be the corresponding value function of (3) for any given player before arrival. Using the constant control $\lambda \equiv 0$ shows that $R(1) \leq \bar{v}(r) \leq R(r)$, and hence $\bar{v}(1-) = R(1)$; recall that R is left continuous at $r = 1$. Let

$$r_0 = \inf\{r : R(r) = R(1)\}.$$

Using $\lambda \equiv 0$ also shows that $\bar{v} \equiv R(1)$ on $[r_0, 1]$, whereas $\bar{v} < R$ on $[0, r_0)$ because the only control with zero cost merely attains $R(1) < R(r)$. The inequality $R(1) < R(r)$ on $[0, r_0)$ also implies that $\bar{\lambda} > 0$ a.e. on $[0, r_0)$, and in fact, recalling that $\bar{\lambda} \in \Lambda$, even that $\bar{\lambda}$ is a.e. uniformly bounded away from zero on any interval $[0, r_1]$ with $r_1 < r_0$.

A control-theoretic argument shows that $\bar{v}$ is Lipschitz on any such interval. Indeed, let $0 \leq r < r_1$ and choose $h > 0$ such that $r + 2h < r_1$. We may compare the optimal control $\bar{\lambda}$ started at r with a control λ such that

$$\lambda = \begin{cases} 0 & \text{on } [r, r+h), \\ a\bar{\lambda} & \text{on } [r+h, r+2h), \\ \bar{\lambda} & \text{on } [r+2h, 1), \end{cases}$$

where the constant $a \geq 1$ is chosen such that the integral of λ along ρ over $[r, r+2h]$ coincides with the integral of $\bar{\lambda}$ over the same interval. Then both the extra cost of λ and the loss of expected reward are bounded by a constant (depending only on r_1) times h because if $\bar{\lambda}$ achieves a rank in $[r+2h, 1)$, then λ achieves the same rank, whereas the probability of $\bar{\lambda}$ ranking in $[r, r+2h)$ is bounded by a constant times h . Because λ is an admissible control for the control problem started at $r+h$, it follows that $0 \leq \bar{v}(r) - \bar{v}(r+h) \leq Ch$ as claimed. In particular, $\bar{v}$ is absolutely continuous and a.e. differentiable.

By dynamic programming, the value function $\bar{v}$ must then a.e. satisfy the Hamilton–Jacobi equation

$$\sup_{l \geq 0} \{l[R(r) - v(r)] - c(r)l^2\} + \bar{\lambda}(r)(1-r)v'(r) = 0 \text{ on } [0, 1], \quad v(1) = R(1).$$

Moreover, the optimal control $\bar{\lambda}$ must attain the supremum a.e. Recalling that $\bar{v} = R$ on $[r_0, 1]$ and $\bar{v} < R$ on $[0, r_0)$, it follows that

$$\bar{\lambda}(r) = \frac{R(r) - \bar{v}(r)}{2c(r)} \quad \text{a.e.} \quad (\text{A.1})$$

and that $\bar{v}$ satisfies

$$R(r) - v(r) + 2(1-r)v'(r) = 0 \text{ a.e. on } [0, r_0), \quad v \equiv R(1) \text{ on } [r_0, 1]. \quad (\text{A.2})$$

(2) Using the regularity of R , we see that the function v of (5) is Lipschitz on $[0, 1]$ and satisfies $v(1-) = R(1)$. We extend v to a Lipschitz function on $[0, 1]$ by setting $v(1) = R(1)$. A direct calculation shows that v satisfies (A.2) and that λ^* of (4) is the corresponding maximizer (A.1). Moreover, $\lambda^* \in \Lambda$. A verification argument then yields that v is the value function and λ^* is an optimal control.

(3) It remains to show that (A.2) has at most one absolutely continuous solution. Indeed, if v_1 and v_2 are solutions, then $w = v_1 - v_2$ is absolutely continuous and satisfies

$$2(1-r)w'(r) = w(r) \text{ a.e. on } [0, r_0), \quad w \equiv 0 \text{ on } [r_0, 1].$$

If $r_0 < 1$, this is a Lipschitz ODE, and it follows directly that $w \equiv 0$ is the unique solution. If $r_0 = 1$, we set $u(t) = w(1 - e^{-2t})$ for $t \geq 0$; then u satisfies $u'(t) = u(t)$ on $[0, \infty)$, and hence $u(t) = u(0)e^t$. But $u(\infty) = w(1) = 0$ then yields that $u \equiv 0$, and thus $w \equiv 0$ as desired. $\square$

Proof of Proposition 1. Let $V(r) = E[R(\rho(\tau)) | \rho(0) = r]$. The ODE for ρ has the unique solution $\rho(t) = 1 - (1-r)e^{-2\lambda_0 t}$ for $t \geq 0$. Hence,

$$V(r) = E[R(1 - (1-r)e^{-2\lambda_0 \tau})] = \int_0^\infty \lambda_0 e^{-\lambda_0 x} R(1 - (1-r)e^{-2\lambda_0 x}) dx.$$

A change of variables then shows that $V(r)$ coincides with (5). $\square$

Proof of Proposition 2. Note that R_n, R have a uniform upper bound given by $\sup_n R_n(0)$. Moreover, by monotonicity, the convergence $R_n \rightarrow R$ is uniform on each interval of continuity of R . It follows directly from (5) that the value functions v_n converge uniformly to their counterpart v . Similarly, (4) yields that the optimal controls λ_n^* converge pointwise to their counterpart λ^* and uniformly on each interval of Lipschitz continuity of R . Moreover, there is a uniform upper bound for the sequence (λ_n^*) . By the ODE (2), this entails an upper bound for the Lipschitz constants of (ρ_n) . Thus, after passing to a subsequence, (ρ_n) converges uniformly to a limit $\bar{\rho}$. To verify that $\bar{\rho} = \rho$, it suffices to show that $\bar{\rho}$ solves the ODE (2) defining ρ on each of the mentioned intervals, and that follows from the uniform convergence of (λ_n^*) . Finally, by a subsequence argument, the entire sequence (ρ_n) must converge to ρ . $\square$

A.2. Proofs for Section 3 As a preparation for the proof of Theorem 2, we first show that no reward should be distributed to ranks below α —this is quite intuitive because the planner does not care about agents arriving after rank α . The converse is also true.

Lemma 2. *The value T_α^* of (8) does not change if the infimum is restricted to $R \in \mathcal{R}$ satisfying $R > 0$ on $[0, \alpha]$ and $R = 0$ on $(\alpha, 1]$.*

Proof. For the first property, suppose that $R \in \mathcal{R}$ vanishes at $r \in [0, \alpha]$ and hence on $[r, 1]$. Then the equilibrium effort λ of (4) also vanishes at r , which by the Kolmogorov equation (2) implies that the state $\rho(t)$ never exceeds r and hence that $T_\alpha(R) = \infty$.

To see the second property, let $R \in \mathcal{R}$, and set $\hat{R} = R\mathbf{1}_{[0, \alpha]}$; then $\hat{R} \in \mathcal{R}$. For $r \in [0, \alpha]$, the corresponding equilibrium efforts λ and $\hat{\lambda}$ of (4) satisfy

$$\hat{\lambda}(r) = \frac{\hat{R}(r) - \frac{1}{2\sqrt{1-r}} \int_r^1 \frac{\hat{R}(y)}{\sqrt{1-y}} dy}{2c(r)} \geq \frac{R(r) - \frac{1}{2\sqrt{1-r}} \int_r^1 \frac{R(y)}{\sqrt{1-y}} dy}{2c(r)} = \lambda(r),$$

and the inequality is strict if $\int_\alpha^1 R(y)/\sqrt{1-y} dy > 0$. As a result, if R does not vanish on $(\alpha, 1]$, then $\hat{R}$ produces a strictly larger equilibrium effort on $[0, \alpha]$, and hence $T_\alpha(\hat{R}) < T_\alpha(R)$ whenever $T_\alpha(R) < \infty$, by (2). $\square$

We can now show the theorem through a calculus-of-variations argument.

Proof of Theorem 2. Let

$$\mathcal{R}' = \{R \in \mathcal{R} : \int R(r) dr \leq B, T_\alpha(R) < \infty, R\mathbf{1}_{(\alpha, 1]} = 0\}.$$

As noted, $\mathcal{R}' \neq \emptyset$ by (6), and by Lemma 2, it is sufficient to show that R^* is the unique optimizer in $\mathcal{R}'$. Moreover, we have seen in the proof of Lemma 2 that λ^* is a.e. strictly positive on $[0, \alpha]$ for $R \in \mathcal{R}'$. Hence, ρ is strictly increasing, and we have $T_r(R) = \rho^{-1}(r)$ for $r \in [0, \alpha]$. Differentiating $\rho^{-1}(r)$ and using the ODE (2) for ρ and (4), we then obtain that

$$T_\alpha(R) = \int_0^\alpha \frac{1}{(1-r)\lambda^*(r)} dr = \int_0^\alpha \frac{2c(r)}{\sqrt{1-r} \left(R(r)\sqrt{1-r} - \int_r^\alpha \frac{R(s)}{2\sqrt{1-s}} ds \right)} dr, \quad R \in \mathcal{R}'. \quad (\text{A.3})$$

We see from (A.3) that $R \mapsto T_\alpha(R)$ is strictly convex on $\mathcal{R}'$, up to a.e. equivalence. This implies that there is at most one optimal $R \in \mathcal{R}'$.

Next, we derive a sufficient condition for optimality. We first reparametrize the optimization problem: for $R \in \mathcal{R}'$, we consider

$$f(r) = f_R(r) = R(r)\sqrt{1-r} - \int_r^\alpha \frac{R(s)}{2\sqrt{1-s}} ds, \quad r \in [0, \alpha].$$

The mapping $R \mapsto f_R$ is one-to-one on $\mathcal{R}'$ because R can be recovered from f via

$$R(r) = \frac{f(r)}{\sqrt{1-r}} + \int_r^\alpha \frac{f(s)}{2(1-s)^{3/2}} ds. \quad (\text{A.4})$$

Indeed, if R is differentiable, then $f'(r) = \sqrt{1-r}R'(r)$, and integration by parts yields

$$R(r) = R(\alpha) - \int_r^\alpha \frac{f'(s)}{\sqrt{1-s}} ds = \frac{f(\alpha)}{\sqrt{1-\alpha}} - \int_r^\alpha \frac{f'(s)}{\sqrt{1-s}} ds = \frac{f(r)}{\sqrt{1-r}} + \int_r^\alpha \frac{f(s)}{2(1-s)^{3/2}} ds,$$

and now the claim for general $R \in \mathcal{R}'$ follows by approximation. Fubini's theorem shows that $\int_0^\alpha R(r) dr = \frac{1}{2} \int_0^\alpha \frac{(2-r)f(r)}{(1-r)^{3/2}} dr$. Thus, recalling that $\alpha < 1$, the image $\mathcal{F}$ of $\mathcal{R}'$ under $R \mapsto f_R$ is the convex set of all piecewise Lipschitz, nonnegative, decreasing functions $f : [0, \alpha] \rightarrow \mathbb{R}$ such that (A.6) is finite, and the budget constraint

$$\frac{1}{2} \int_0^\alpha \frac{(2-r)f(r)}{(1-r)^{3/2}} dr \leq B \quad (\text{A.5})$$

is satisfied. We write $T_\alpha(f_R)$ for $T_\alpha(R)$ by a slight abuse of notation; then by (A.3) we have

$$T_\alpha(f) = \int_0^\alpha \frac{2c(r)}{\sqrt{1-r}f(r)} dr, \quad f \in \mathcal{F}. \quad (\text{A.6})$$

The mapping $f \mapsto T_\alpha(f)$ is convex and finite valued, and clearly, $f^* \in \mathcal{F}$ is optimal if and only if

$$\varepsilon \mapsto \phi(\varepsilon) = T((1-\varepsilon)f^* + \varepsilon f), \quad [0, 1] \rightarrow \mathbb{R}$$

attains a minimum at $\varepsilon = 0$ for all $f \in \mathcal{F}$. By convexity, this function has a right derivative $\phi'(0)$ at $\varepsilon = 0$, and ϕ attains a minimum at $\varepsilon = 0$ if and only if $\phi'(0) \geq 0$. Note that for any convex function φ , the right difference quotient $(\varphi(x + \varepsilon) - \varphi(x))/\varepsilon$ satisfies $\varphi'(x) \leq (\varphi(x + \varepsilon) - \varphi(x))/\varepsilon \leq \varphi(x + 1) - \varphi(x)$ for $\varepsilon \leq 1$. Using these bounds and dominated convergence, we see that $\phi'(0)$ can be computed by differentiation under the integral:

$$\phi'(0) = \int_0^\alpha \frac{-2c(r)(f(r) - f^*(r))}{\sqrt{1 - rf^*(r)^2}} dr. \quad (\text{A.7})$$

Let R^* be as in (10); then the corresponding function $f^* = f_{R^*}$ is given by

$$f^*(r) = \frac{B}{C} \sqrt{\frac{c(r)(1-r)}{2-r}}, \quad r \in [0, \alpha]. \quad (\text{A.8})$$

Recalling that $\alpha < 1$ and that $c(r)(1-r)/(2-r)$ is decreasing, we verify directly that $f^* \in \mathcal{F}$. Moreover, the budget constraint (A.5) is satisfied with equality.

Fix an arbitrary $f \in \mathcal{F}$. Using the expression (A.8) in the denominator of (A.7), we have

$$\phi'(0) = \frac{-2C^2}{B^2} \int_0^\alpha \frac{(2-r)(f(r) - f^*(r))}{(1-r)^{3/2}} dr.$$

Because f^* satisfies (A.5) with equality and f satisfies the same with inequality, the above integral is nonpositive. Thus, $\phi'(0) \geq 0$, showing that $f^* \in \mathcal{F}$ is optimal and hence that $R^* \in \mathcal{R}$ is optimal. The formula (11) for T_α^* is then obtained from (A.6) and (A.8), and the formula for λ^* follows from (4). $\square$

A.3. Proofs for Section 4

Proof of Proposition 3. Fix player i and suppose that all other players use a control λ_{-i} . By dynamic programming, the value function v_n of player i before arrival satisfies

$$\sup_{\lambda_i \geq 0} \{ \lambda_i [R_{n+1} - v_n] - c_n \lambda_i^2 \} + \lambda_{-i}(n)(N - n - 1)(v_{n+1} - v_n) = 0$$

for $0 \leq n \leq N - 2$. For $n = N - 1$, the same holds by Remark 3 and our convention that $v_N = R_N$. Thus, (13) is the optimal control for player i , and it follows that

$$\frac{(R_{n+1} - v_n)^2}{4c_n} + \lambda_{-i}(n)(N - n - 1)(v_{n+1} - v_n) = 0.$$

Assuming inductively that $v_{n+1} \leq R_{n+1}$, this quadratic equation has a unique nonnegative root v_n , and v_n satisfies $0 \leq v_n \leq R_{n+1} \leq R_n$. In a given equilibrium, the consistency condition $\lambda_i = \lambda_{-i}$ implies that

$$R_{n+1} - v_n + 2(N - n - 1)(v_{n+1} - v_n) = 0, \quad n = 0, \dots, N - 1, \quad (\text{A.9})$$

or, equivalently, (12), which clearly has a unique solution. Conversely, we can verify directly that (12) and (13) define an equilibrium. $\square$

Proof of Theorem 3. We first observe that T_{n_0} is a sum of independent exponential random variables (whenever finite), and thus

$$ET_{n_0} = \sum_{n=0}^{n_0-1} \frac{1}{(N - n)\lambda_n}. \quad (\text{A.10})$$

Moreover, similarly as in Lemma 2, it suffices to consider reward schemes with $R_n > 0$ for $n \leq n_0$ and $R_n = 0$ for $n > n_0$. Indeed, the first claim is immediate from (13). To obtain the second, we argue by contradiction and compare (R_n) with the scheme defined by $\hat{R}_n = R_n \mathbf{1}_{\{n \leq n_0\}}$. Proposition 3 implies that $\hat{R}_n$ leads to a strictly larger equilibrium control before n_0 and hence a strictly smaller completion time. Thus, we only consider reward schemes up to $n = n_0$ in what follows.

From (12), (13), and (A.10), we see that $ET_{n_0} : \mathbb{R}_+^{n_0} \rightarrow [0, \infty]$ is a strictly convex, continuous function of $(R_1, \dots, R_{n_0})$. Moreover, the feasible set defined by $R_1 \geq R_2 \geq \dots \geq R_{n_0} \geq 0$ and $\sum_{n=1}^N R_n \leq NB$ is nonempty, convex, and compact. As a result, there exists a unique optimal reward scheme. In the remainder of the proof, we determine this reward scheme explicitly. To that end, define $x_{n_0} = 0$ and $x_n = R_{n+1} - v_n$ for $n < n_0$. By (13) and (A.10), the objective function can then be expressed as

$$ET_{n_0} = \sum_{n=0}^{n_0-1} \frac{2c_n}{(N - n)x_n}.$$

From (12) we obtain that $R_{n+1} - R_{n+2} = \frac{1+2(N-n-1)}{2(N-n-1)}x_n - x_{n+1}$, and thus the total reward $\sum_{n=1}^{n_0} R_n = \sum_{n=1}^{n_0} n(R_n - R_{n+1})$ can be expressed as

$$\sum_{n=1}^{n_0} R_n = \sum_{n=1}^{n_0} n \left\{ \frac{1+2(N-n)}{2(N-n)} x_{n-1} - x_n \right\} = \sum_{n=0}^{n_0-1} \frac{2N-n-1}{2(N-n-1)} x_n.$$

The constrained optimization problem of finding $x_0, \dots, x_{n_0-1} \in \mathbb{R}$ that

$$\text{minimize } \sum_{n=0}^{n_0-1} \frac{2c_n}{(N-n)x_n} \quad \text{subject to } \sum_{n=0}^{n_0-1} \frac{2N-n-1}{2(N-n-1)} x_n \leq NB,$$

can be solved using the method of Lagrange multipliers. One finds that the optimal x_n are given by

$$x_n = 2 \sqrt{\frac{c_n(N-n-1)}{\theta(N-n)(2N-n-1)}}, \quad n = 0, \dots, n_0-1,$$

where the Lagrange multiplier θ satisfies

$$\sqrt{\theta} = \frac{1}{NB} \sum_{n=0}^{n_0-1} \sqrt{\frac{c_n(2N-n-1)}{(N-n)(N-n-1)}}.$$

Note that θ and x_n are related to C and y_n of Theorem 3 by $C = B\sqrt{N\theta}/2$ and $x_n = 2y_n/\sqrt{N\theta} = By_n/C$. We have that $R_{n_0} \geq 0$ as $x_{n_0-1} \geq 0$, and $R_n - R_{n+1} = \frac{B}{C} \left(\frac{1+2(N-n)}{2(N-n)} y_{n-1} - y_n \right) \geq 0$ for $n < n_0$; here the last inequality is equivalent to (14). Thus, the reward scheme (R_n^*) associated with the optimizer (x_n) is indeed nonnegative and decreasing, and it follows that (R_n^*) is the optimal reward scheme. The formulas for R^* , $T_{n_0}^*$, and λ^* follow by direct calculation. $\square$

A.4. Proofs for Section 5

Proof of Theorem 4. Let $N \geq 4$ be large enough such that $\delta := 1/N$ satisfies $r_m < 1 - \sqrt{\delta} - \delta$. This implies, in particular, that $1 - r_m > \delta + \sqrt{\delta} > 2\delta$. We may rewrite the recursion (12) for v_n^N as

$$v_n^N = g\left(\frac{n}{N}\right) v_{n+1}^N + \left(1 - g\left(\frac{n}{N}\right)\right) R_{n+1}^N, \quad (\text{A.11})$$

where

$$g(r) = \frac{2(1-r-\delta)}{\delta + 2(1-r-\delta)}.$$

Using (5), we write the mean field value in a similar form:

$$v\left(\frac{n}{N}\right) = f\left(\frac{n}{N}\right) v\left(\frac{n+1}{N}\right) + \left(1 - f\left(\frac{n}{N}\right)\right) R_{n+1}^N + E_{1,n}, \quad (\text{A.12})$$

where

$$f(r) := \sqrt{\frac{1-r-\delta}{1-r}} \quad \text{and} \quad E_{1,n} := \frac{1}{2\sqrt{1-n\delta}} \int_{n\delta}^{(n+1)\delta} \frac{R(y) - R_{[yN]}^N}{\sqrt{1-y}} dy.$$

Next, we estimate $\Delta_n := |v_n^N - v(\frac{n}{N})|$. For the last rank, we have $\Delta_N = |R_N^N - R(1)| \leq K\delta$ by (15). For the second-to-last rank, because $1 - \delta \geq 1 - \sqrt{\delta} > r_m + \delta$, (15) again implies that

$$\Delta_{N-1} = |R_{N-1}^N - v(1-\delta)| \leq \frac{1}{2\sqrt{\delta}} \int_{1-\delta}^1 \frac{|R_{[yN]}^N - R(y)|}{\sqrt{1-y}} dy \leq K\delta.$$

For $n \leq N-2$, subtracting (A.12) from (A.11) and using that $g(r) \leq 1$ when $0 \leq r \leq 1 - \delta$, we obtain that

$$\Delta_n \leq \Delta_{n+1} + E_{1,n} + E_{2,n}, \quad (\text{A.13})$$

where

$$E_{2,n} = \left| g\left(\frac{n}{N}\right) - f\left(\frac{n}{N}\right) \right| \cdot \left| v\left(\frac{n+1}{N}\right) - R_{n+1}^N \right|.$$

Next, we estimate $E_{1,n}$ and $E_{2,n}$. To that end, it will be useful to define $n_0 := N - \lceil \sqrt{N} \rceil$ and note that $n \leq n_0$ if and only if $n\delta \leq 1 - \sqrt{\delta}$.

Estimation of $E_{1,n}$. If $n_0 + 1 \leq n \leq N - 2$, then $r_m + \delta < n\delta \leq 1 - 2\delta$, and (15) implies

$$E_{1,n} \leq K\delta \left(1 - \sqrt{\frac{1 - n\delta - \delta}{1 - n\delta}} \right) \leq K\delta.$$

If $0 \leq n \leq n_0$ and $[n\delta, (n+1)\delta] \subseteq I_i^N$ for some i , then (15) and $1 - n\delta \geq \sqrt{\delta}$ imply

$$E_{1,n} \leq K\delta \frac{\frac{\delta}{1-n\delta}}{1 + \sqrt{\frac{1-n\delta-\delta}{1-n\delta}}} \leq K\delta^{3/2}.$$

We observe that there are at most $3m$ indices n for which $n \leq n_0$, and $[n\delta, (n+1)\delta]$ is not contained in any I_i^N . For these n , we use the boundedness of R and the inequality $1 - n\delta \geq 1 - r_m - \delta > (1 - r_m)/2$ to obtain

$$E_{1,n} \leq (R(0) + K\delta) \left(1 - \sqrt{\frac{1 - n\delta - \delta}{1 - n\delta}} \right) < \frac{2(R(0) + K/4)\delta}{(1 - r_m)} =: C_0\delta.$$

In summary,

$$E_{1,n} \leq \begin{cases} K\delta, & \text{if } n_0 + 1 \leq n \leq N - 2, \\ K\delta^{3/2}, & \text{if } n \leq n_0 \text{ and } [n\delta, (n+1)\delta] \subseteq I_i^N \text{ for some } i, \\ C_0\delta, & \text{otherwise (at most } 3m \text{ instances).} \end{cases} \quad (\text{A.14})$$

Estimation of $E_{2,n}$. Taylor's theorem implies that for all $x \in [0, 1]$,

$$0 \leq \sqrt{x} - \frac{2x}{1+x} \leq \frac{1}{2}(1-x)^2 \left(\frac{1}{4x^{3/2}} + \frac{4}{(1+x)^3} \right).$$

Using this with $x = 1 - \frac{\delta}{1-r}$ yields

$$|g(r) - f(r)| \leq \frac{\delta^2}{(1-r)^2} h \left(1 - \frac{\delta}{1-r} \right),$$

where $h(x) := \frac{1}{8x^{3/2}} + \frac{2}{(1+x)^3}$. For $0 \leq n \leq Nr_m$, we have $\frac{\delta}{1-n\delta} \leq \frac{\delta}{1-r_m} \leq \frac{1}{2}$, and thus

$$E_{2,n} \leq (R(0) + K/4) \frac{\delta^2}{(1-r_m)^2} h \left(\frac{1}{2} \right) =: C_1\delta^2.$$

For $Nr_m < n \leq N - 2$, we have $(n+1)\delta > r_m + \delta$, and it follows that

$$\begin{aligned} \left| v \left(\frac{n+1}{N} \right) - R_{n+1}^N \right| &= \left| \frac{1}{2\sqrt{1 - (n+1)\delta}} \int_{(n+1)\delta}^1 \frac{R(y) - R((n+1)\delta)}{\sqrt{1-y}} dy + R((n+1)\delta) - R_{n+1}^N \right| \\ &\leq K(1 - (n+1)\delta) + K\delta = K(1 - n\delta). \end{aligned}$$

Consequently,

$$E_{2,n} \leq K(1 - n\delta) \frac{\delta^2}{(1 - n\delta)^2} h \left(1 - \frac{\delta}{1 - n\delta} \right) \leq \frac{K}{2} h \left(\frac{1}{2} \right) \delta < \frac{K}{2} \delta,$$

where we have used that $1 - n\delta \geq 2\delta$.

As in the estimation of $E_{1,n}$, the bound can be improved if $Nr_m < n \leq n_0$, and thus $1 - n\delta \geq \sqrt{\delta}$. In this case, we have

$$E_{2,n} \leq Kh \left(\frac{1}{2} \right) \delta^{3/2} < K\delta^{3/2}.$$

In summary,

$$E_{2,n} \leq \begin{cases} \frac{1}{2}K\delta, & \text{if } n_0 + 1 \leq n \leq N - 2, \\ K\delta^{3/2}, & \text{if } Nr_m < n \leq n_0, \\ C_1\delta^2, & \text{if } n \leq Nr_m. \end{cases} \quad (\text{A.15})$$

Combining the Estimates. Combining (A.13), (A.14), and (A.15) and recalling that $\Delta_{N-1} \leq K\delta$, we obtain for $n_0 + 1 \leq n \leq N - 2$ that

$$\Delta_n \leq \Delta_{N-1} + \frac{3}{2}K\delta(N-1-n) \leq K\delta + \frac{3}{2}K\sqrt{\delta} < \frac{5}{2}K\sqrt{\delta}$$

and for $n \leq n_0$ that

$$\begin{aligned} \Delta_n &\leq \Delta_{n_0+1} + (2K + C_1)\delta^{3/2}(n_0 + 1 - n) + 3mC_0\delta \\ &\leq \frac{5}{2}K\sqrt{\delta} + (2K + C_1)\sqrt{\delta}(1 - \sqrt{\delta} + \delta) + 3mC_0\delta < C_2\sqrt{\delta}. \end{aligned}$$

Putting everything together, we have $\sup_{0 \leq n \leq N} \Delta_n \leq \max(5K/2, C_2)/\sqrt{N}$. It remains to note that

$$|v'(r)| = \frac{|R(r) - v(r)|}{2(1-r)} \leq \begin{cases} \frac{K}{2}, & r_m < r < 1, \\ \frac{R(0)}{2(1-r_m)}, & 0 \leq r \leq r_m, \end{cases}$$

and consequently,

$$\left| v_{\lfloor rN \rfloor}^N - v(r) \right| \leq \Delta_{\lfloor rN \rfloor} + \left| v\left(\frac{\lfloor rN \rfloor}{N}\right) - v(r) \right| \leq \Delta_{\lfloor rN \rfloor} + \frac{\|v'\|_\infty}{N} < \frac{C_3}{\sqrt{N}}.$$

Because $\lambda^*(r) = \frac{R(r)-v(r)}{2c(r)}$ and $\lambda_n^N = \frac{R_n^N - v_n^N}{2c_n^N}$, the convergence of λ^N follows from the uniform convergence of the value functions and the cost coefficients, the almost uniform convergence of the reward scheme, and the uniform boundedness of $1/c$. $\square$

Proof of Theorem 5. We first observe the convergence of the Riemann sum

$$C^N := \frac{1}{2N} \sum_{n=0}^{\lfloor \alpha N \rfloor - 1} \sqrt{\frac{c_n^N(2 - \frac{n+1}{N})}{(1 - \frac{n}{N})(1 - \frac{n+1}{N})}} \rightarrow \frac{1}{2} \int_0^\alpha \frac{\sqrt{c(r)(2-r)}}{(1-r)} dr =: C.$$

The rate of convergence is $O(1/N)$ because c^N converges to c (which is bounded away from zero) uniformly at rate $O(1/N)$ and $\sqrt{2-r}/(1-r)$ is Lipschitz continuous on $[0, \alpha]$. Using the formulas of Theorems 2 and 3, we conclude that

$$\lim_{N \rightarrow \infty} ET_{\lfloor \alpha N \rfloor}^N = \lim_{N \rightarrow \infty} \frac{4(C^N)^2}{B} = \frac{4C^2}{B} = T_\alpha^*$$

with rate $O(1/N)$. Turning to the convergence of the reward schemes, we similarly observe that for $r \leq \alpha$,

$$y_{\lfloor rN \rfloor - 1} = \sqrt{\frac{c_{\lfloor rN \rfloor - 1}^N N(N - \lfloor rN \rfloor)}{(N - \lfloor rN \rfloor + 1)(2N - \lfloor rN \rfloor)}} \rightarrow \sqrt{\frac{c(r)}{2-r}} = y(r)$$

and

$$\frac{1}{N} \sum_{k=\lfloor rN \rfloor - 1}^{\lfloor \alpha N \rfloor - 1} \frac{y_k}{1 - \frac{k+1}{N}} \rightarrow \int_r^\alpha \frac{y(s)}{1-s} ds,$$

uniformly in $r \in [0, \alpha]$ with rate $O(1/N)$. Hence, the formulas in Theorems 2 and 3 yield that

$$R_{\lfloor rN \rfloor}^N = \frac{B}{C^N} \left\{ y_{\lfloor rN \rfloor - 1} + \frac{1}{2} \sum_{k=\lfloor rN \rfloor - 1}^{\lfloor \alpha N \rfloor - 1} \frac{y_k}{N - k - 1} \right\} \rightarrow \frac{B}{C} \left\{ y(r) + \frac{1}{2} \int_r^\alpha \frac{y(s)}{1-s} ds \right\} = R^*(r)$$

uniformly in $r \in [0, \alpha]$ with rate $O(1/N)$. $\square$

Proof of Corollary 2. Theorem 4 and Remark 4 imply that the N -player equilibrium control $\lambda_{\lfloor rN \rfloor}^{(N)}$ for $R^{(N)}$ converges uniformly to the mean field equilibrium control $\lambda^*(r)$ for R^* at rate $O(1/N)$. It follows that

$$f_N(r) := (1 - \lfloor rN \rfloor) \lambda_{\lfloor rN \rfloor}^{(N)} \rightarrow (1-r) \lambda^*(r) =: f(r)$$

uniformly at rate $O(1/N)$. Because λ^* is uniformly bounded away from zero on $[0, \alpha]$, the same holds for f and f_N with large N . Thus, using (A.10) and (A.3),

$$\left| ET_{\lfloor \alpha N \rfloor}^{(N)} - T_\alpha^* \right| = \left| \sum_{n=0}^{\lfloor \alpha N \rfloor - 1} \frac{1}{(N-n) \lambda_n^{(N)}} - \int_0^\alpha \frac{dr}{(1-r) \lambda^*(r)} \right| \leq \int_0^{\lfloor \alpha N \rfloor} \frac{|f(r) - f_N(r)|}{f_N(r) f(r)} dr = O(1/N).$$

Because Theorem 5 shows that $|T_\alpha^* - ET_{\lfloor \alpha N \rfloor}^N| = O(1/N)$, the claim follows. $\square$

Appendix B. Exact Law of Large Numbers

In this section, we detail a setting such that the exact law of large numbers holds for a continuum of essentially pairwise independent random variables. Let $(I, \mathcal{F}, \mu)$ be an atomless (hence uncountable) probability space, and let $(\Omega, \mathcal{F}, P)$ be another probability space.

Definition 1. A family $(f_i)_{i \in I}$ of random variables on $(\Omega, \mathcal{F}, P)$ is *essentially pairwise independent* if for μ -almost all $i \in I$, f_i is independent of f_j for μ -almost all $j \in I$. The family is *essentially pairwise independent and identically distributed* (i.i.d.) if, in addition, all f_i have the same distribution.

In what follows, we need to work on a probability space that is larger than the usual product⁶ $(I \times \Omega, \mathcal{F} \otimes \mathcal{F}, \mu \otimes P)$ because the latter does not support relevant families of i.i.d. random variables (see, e.g., Sun [38, proposition 2.1]). Following Sun [38], a probability space $(I \times \Omega, \Sigma, \nu)$ is an *extension* of the product $(I \times \Omega, \mathcal{F} \otimes \mathcal{F}, \mu \otimes P)$ if Σ contains $\mathcal{F} \otimes \mathcal{F}$ and the restriction of ν to $\mathcal{F} \otimes \mathcal{F}$ coincides with $\mu \otimes P$. It is a *Fubini extension* if, in addition, any ν -integrable⁷ function $f : I \times \Omega \rightarrow \mathbb{R}$ satisfies the assertion of Fubini's theorem; that is,

- (1) For μ -almost all $i \in I$, the function $f(i, \cdot)$ is P -integrable;
- (2) For P -almost all $\omega \in \Omega$, the function $f(\cdot, \omega)$ is μ -integrable; and
- (3) $i \mapsto \int f(i, \cdot) dP$ is μ -integrable, $\omega \mapsto \int f(\cdot, \omega) d\mu$ is P -integrable, and

$$\int f d\nu = \int \int f(i, \omega) P(d\omega) \mu(di) = \int \int f(i, \omega) \mu(di) P(d\omega).$$

Let $(I \times \Omega, \Sigma, \nu)$ be a Fubini extension of $(I \times \Omega, \mathcal{F} \otimes \mathcal{F}, \mu \otimes P)$. Then essentially pairwise independent families satisfy an exact version of the law of large numbers. The following is a special case of Sun [38, corollary 2.9].

Proposition 4 (Exact Law of Large Numbers). *Let $f : I \times \Omega \rightarrow \mathbb{R}$ be ν -integrable. If $f(i, \cdot)$, $i \in I$ are essentially pairwise i.i.d. with a distribution having mean m , then $\int f(\cdot, \omega) d\mu = m$ for P -almost all $\omega \in \Omega$.*

Next, we turn to the existence of such a setting. The space $(I \times \Omega, \Sigma, \nu)$ is called *rich* if there exists a Σ -measurable function $f : I \times \Omega \rightarrow \mathbb{R}$ such that $f(i, \cdot)$, $i \in I$ are essentially pairwise i.i.d. with a uniform distribution on $[0, 1]$. By composing f with a suitable function, it then follows that $(I \times \Omega, \Sigma, \nu)$ supports essentially pairwise i.i.d. families with any given distribution.

Lemma 3. *There exist atomless probability spaces $(I, \mathcal{F}, \mu)$ and $(\Omega, \mathcal{F}, P)$ such that $(I \times \Omega, \mathcal{F} \otimes \mathcal{F}, \mu \otimes P)$ admits a rich Fubini extension.*

This is part of the assertion of Sun [38, proposition 5.6], which also shows that one can take $I = [0, 1]$ and $\Omega = \mathbb{R}^{[0,1]}$. The main result of Sun and Zhang [39] shows that, in addition, one can take μ to be an extension of the Lebesgue measure (but not the Lebesgue measure itself). A different construction is presented by Podczek [33].

Endnotes

¹“Decreasing” and “increasing” are understood in the nonstrict sense throughout the paper.

²While feedback controls seem to be the most natural choice for our model, an alternate formulation may use open-loop controls, meaning that agents choose functions of time rather than state p . Because in our game a single agent has negligible impact on the state process, this would lead to the same equilibrium as for feedback controls.

³That is, $[0, 1]$ is the union of finitely many intervals on which λ is Lipschitz.

⁴A normalization of the budget is necessary for a convergence result as in the subsequent section. We have done that implicitly by seeing B as the total budget in the mean field limit and as the per capita budget in the N -player setting. Equivalently, one could normalize the mass of the population by assigning mass $1/N$ to each agent in the N -player game and see B as the total budget.

⁵This covers, for example, $c_n^N = c(n/N)$. We do not consider a more general convergence as in (15) because that would lead to more complicated statements below without substantially enhancing the scope.

⁶Here we use the convention that the product σ -field $\mathcal{F} \otimes \mathcal{F}$ is completed.

⁷That is, f is measurable for the ν -completion of Σ and $\int |f| d\nu < \infty$.

References

- [1] Bayraktar E, Cohen A (2018) Analysis of a finite state many player game using its master equation. *SIAM J. Control Optim.* 56(5):3538–3568.
- [2] Bayraktar E, Zhang Y (2016) A rank-based mean field game in the strong formulation. *Electronic Comm. Probab.* 21(72):1–12.
- [3] Bensoussan A, Chau MHM, Yam SCP (2016) Mean field games with a dominating player. *Appl. Math. Optim.* 74(1):91–128.
- [4] Bensoussan A, Frehse J, Yam SCP (2013) *Mean Field Games and Mean Field Type Control Theory*, Springer Briefs in Mathematics (Springer, New York).
- [5] Bertucci C (2018) Optimal stopping in mean field games, an obstacle problem approach. *J. Math. Pures Appl.* 120:165–194.
- [6] Cao D (2014) Racing under uncertainty: Boundary value problem approach. *J. Econom. Theory* 151:508–527.
- [7] Cardaliaguet P, Delarue F, Lasry JM, Lions PL (2015) The master equation and the convergence problem in mean field games. Working paper, Université Paris-Dauphine, Paris.
- [8] Carmona R, Delarue F (2018) *Probabilistic Theory of Mean Field Games with Applications I* (Springer International, Cham Switzerland).
- [9] Carmona R, Delarue F (2018) *Probabilistic Theory of Mean Field Games with Applications II* (Springer International, Cham Switzerland).

- [10] Carmona R, Wang P (2016) Finite state mean field games with major and minor players. Working paper, Princeton University, Princeton, NJ.
- [11] Carmona R, Delarue F, Lacker D (2017) Mean field games of timing and models for bank runs. *Appl. Math. Optim.* 76(1):217–260.
- [12] Elie R, Possamai D (2016) Contracting theory with competitive interacting agents. Working paper, Université Paris-Est Marne-la-Vallée, Champs-sur-Marne, France.
- [13] Elie R, Mastrolia T, Possamai D (2018) A tale of a principal and many, many agents. *Math. Oper. Res.* 44(2):440–467.
- [14] Feng H, Hobson D (2016) Gambling in contests with random initial law. *Ann. Appl. Probab.* 26(1):186–215.
- [15] Feng H, Hobson D (2016) Gambling in contests with regret. *Math. Finance* 26(3):674–695.
- [16] Gomes DA, Mohr J, Souza RR (2013) Continuous time finite state mean field games. *Appl. Math. Optim.* 68(1):99–143.
- [17] Grossman GM, Shapiro C (1987) Dynamic R&D competition. *Econom. J.* 97(386):372–387.
- [18] Guéant O, Lasry JM, Lions PL (2011) Mean field games and applications. *Paris-Princeton Lectures on Mathematical Finance 2010*, Lecture Notes in Mathematics, vol. 2003 (Springer, Berlin), 205–266.
- [19] Harris C, Vickers J (1987) Racing with uncertainty. *Rev. Econom. Stud.* 54(1):1–21.
- [20] Hörner J (2004) A perpetual race to stay ahead. *Rev. Econom. Stud.* 71(4):1065–1088.
- [21] Huang M (2010) Large-population LQG games involving a major player: The Nash certainty equivalence principle. *SIAM J. Control Optim.* 48(5):3318–3353.
- [22] Huang M, Caines PE, Malhamé RP (2007) Large-population cost-coupled LQG problems with nonuniform agents: Individual-mass behavior and decentralized ϵ -Nash equilibria. *IEEE Trans. Automatic Control* 52(9):1560–1571.
- [23] Huang M, Malhamé RP, Caines PE (2006) Large population stochastic dynamic games: Closed-loop McKean-Vlasov systems and the Nash certainty equivalence principle. *Comm. Inform. Systems* 6(3):221–251.
- [24] Koo HK, Shim G, Sung J (2008) Optimal multi-agent performance measures for team contracts. *Math. Finance* 18(4):649–667.
- [25] Lasry JM, Lions PL (2006) Jeux à champ moyen. I. Le cas stationnaire. *C. R. Math. Acad. Sci. Paris* 343(9):619–625.
- [26] Lasry JM, Lions PL (2006) Jeux à champ moyen. II. Horizon fini et contrôle optimal. *C. R. Math. Acad. Sci. Paris* 343(10):679–684.
- [27] Lasry JM, Lions PL (2007) Mean field games. *Japanese J. Math.* 2(1):229–260.
- [28] Lazear E, Rosen S (1981) Rank-order tournaments as optimum labor contracts. *J. Political Econom.* 89(5):841–864.
- [29] Malueg DA, Tsutsui SO (1997) Dynamic R&D competition with learning. *RAND J. Econom.* 28(4):751–772.
- [30] Moscarini G, Smith L (2011) Optimal dynamic contests. Working paper, Yale University, New Haven, CT.
- [31] Nadtochiy S, Shkolnikov M (2019) Particle systems with singular interaction through hitting times: Application in systemic risk modeling. *Ann. Appl. Probab.* 29(1):89–129.
- [32] Nutz M (2018) A mean field game of optimal stopping. *SIAM J. Control Optim.* 56(2):1206–1221.
- [33] Podczeck K (2010) On existence of rich Fubini extensions. *Econom. Theory* 45(1–2):1–22.
- [34] Reinganum JF (1982) A dynamic game of R and D: Patent protection and competitive behavior. *Econometrica* 50(3):671–688.
- [35] Seel C, Strack P (2013) Gambling in contests. *J. Econom. Theory* 148(5):2033–2048.
- [36] Seel C, Strack P (2016) Continuous time contests with private information. *Math. Oper. Res.* 41(3):1093–1107.
- [37] Shkolnikov M (2012) Large systems of diffusions interacting through their ranks. *Stochastic Processes Their Appl.* 122(4):1730–1747.
- [38] Sun Y (2006) The exact law of large numbers via Fubini extension and characterization of insurable risks. *J. Econom. Theory* 126(1):31–69.
- [39] Sun Y, Zhang Y (2009) Individual risk and Lebesgue extension without aggregate uncertainty. *J. Econom. Theory* 144(1):432–443.
- [40] Taylor CR (1995) Digging for golden carrots: An analysis of research tournaments. *Amer. Econom. Rev.* 85(4):872–890.
- [41] Vojnović M (2015) *Contest Theory* (Cambridge University Press, Cambridge, UK).