

OPTIMAL DYNAMIC CONTESTS*

Giuseppe Moscarini
Yale University[†] and NBER

Lones Smith
University of Michigan[‡]

January 2007
(in progress)

Abstract

We study the design of optimal incentives in a two-player dynamic contest. Two players continuously spend costly effort to attain a score lead, which is also affected by noise. The first player to reach a predetermined score difference (finish line) wins a prize. We focus on the choice of the optimal prize for the winner and on the optimal scoring rule, which penalizes or boosts the leader at each point in the game. We solve for a Symmetric Markov Perfect Equilibrium of the contest, and use it to evaluate a few possible principal's objectives. We find that equilibrium effort is always positive, increasing or hump-shaped in own lead, and the leader always exerts more effort than the follower. These results replicate in our continuous time, continuous state space setting those obtained by Harris and Vickers (1987) in a different discrete time, discrete state model. Our model is more tractable and affords our main innovation, the normative analysis of this game. Due to the strategic interaction, the optimal prize that maximizes expected total agents' output is finite even if effort costs and the value of the prize are of no concern to the principal. Too large a prize and the strategic complementarities of agents' efforts generate an initial war phase, followed by low effort thereafter and whenever the lead is not very small. We conjecture that the optimal scoring rule entails penalizing the leader to keep the laggard from giving up, despite the adverse ex ante incentives of this rule. We show how to solve numerically for the optimal scoring rule.

*This is based on our joint work proposed in a funded National Science Foundation grant (2000–2003) of Lones Smith. Both Giuseppe and Lones acknowledge a joint current NSF grant (2006–9). We thank Daron Acemoglu for useful comments on this preliminary draft.

[†]Department of Economics and Cowles Foundation for Research in Economics, PO Box 208268, New Haven, CT 06520-8268, U.S.A., giuseppe.moscarini@yale.edu, www.econ.yale.edu/~mosca/mosca.html

[‡]Department of Economics, Ann Arbor, MI 48109-1220; www.umich.edu/~lones, lones@umich.edu

1 Introduction

When an agent's performance is curbed by moral hazard, how should a principal optimally design incentives? This question has been analyzed extensively in the single-agent case both in the static and in the repeated contexts, as well as, in a static context, in a multi-agent framework. In the latter case, an important dimension of the problem is team production. If agents' efforts are not separable in the technology to produce the output that the principal cares about, then the principal gains additional latitude to provide incentives. But when technologies are independent, so each agent produces a performance that depends only on her own effort (and noise), should incentives also be independent across agents?

The literature on tournaments suggests a negative answer: setting up a direct competition among agents whose technologies are separate and independent can improve incentives. This literature is almost exclusively concerned with static situations. In reality, however, tournaments are almost always dynamic. Prominent examples are sport leagues and promotion policies within professional companies. The analysis of dynamic tournaments is scarce. The best known example is the race set up by Harris and Vickers (1987) either in one dimension, where only the difference in performances matter, or in two dimensions, where the scale of individual performances also matters to determine rewards. The existing analysis is exclusively positive. We know nothing about the optimal design of a dynamic tournament.

More generally, we know almost nothing about the design of optimal incentives for multiple players who interact repeatedly. How are the rules of a professional sports league optimally designed? Notice that the space of mechanisms is very wide. Even fully describing it appears to be a daunting task. The principal can choose handicaps, length of the tournament, monetary prizes or penalties at each stage and in the whole tournament, rules for the stage game as a function of the current state of the overall game, etc.

In this paper, we attempt a first attack to the problem by focusing attention on a restricted, yet natural and still fairly flexible class of mechanisms in a two-agent game. No rewards and penalties are paid during the game. Payoffs come only at the end. Players continuously exert costly efforts that produce outputs with some noise. The winner is the player who achieves first a pre-determined output difference. I.e. the winner is the first player who outperforms the opponent by a certain margin. The winner gets a prize and the

loser gets nothing, having spent effort costs to no avail. An example is tennis without a tie-breaker. The game can last for ever, although in equilibrium it must always end in finite time a.s. if the game is designed to maximize incentives. While there is a vast literature on the war of attrition, which ends with a one-way concession, we think of a more general strategic situation that “allows for a back and forth.” The lead can pass between opponents many times before the game ends when one player reaches a set finish line. In the specific linear-quadratic example of this paper, a player produces effort at quadratic cost to control linearly the drift of a Brownian Motion, which is his cumulative output. The difference in cumulative outputs is also a Brownian Motion. Payoffs are undiscounted.

The winner is the first player to drive this difference to a finish line in his favor. In the unique Symmetric Markov Perfect Equilibrium of this game, the equilibrium strategy has three properties: optimal effort is always strictly positive, it is either increasing or hump-shaped in own score, and whoever happens to be the leader always exerts pointwise more effort than the follower. So, the initial lead encourages the early leader and discourages the laggard, and the leader slacks off only when his lead becomes sufficiently large and the prize is sufficiently high, but never as much as the laggard. These results replicate exactly those obtained by Harris and Vickers (1987) in their one-dimensional tug-of-war, which is a discrete time, discrete state, exponential-chance version of this game. This suggests that these equilibrium properties transcend the specific setup.

We then turn to our main innovation, the mechanism design problem. We consider a principal who chooses the finish line, the prize, and the instantaneous weighting of the performance difference to maximize expected total output, either gross or net of effort costs. First, we focus on the choice of the finish line and of the optimal prize when the “score” to determine victory is in “natural scale,” namely the difference in cumulative outputs. We find that the optimal finish line is unbounded. This appears to be an artifact of the linear-quadratic setup. Since this affords much tractability, we maintain it and focus on the choice of the optimal prize for exogenously given finish line. Our main conjecture, verified in numerical examples, is that the optimal prize is finite even if the principal does not care either about the agents’ effort cost or the value of the prize. In contrast, under the same assumptions about the principal’s preferences, when facing a single-agent repeated moral hazard problem the optimal payment is obviously unbounded above.

Some numerical examples suggest the reason for this result. When the prize is very high, players engage in an initial “war” phase, where their efforts are very large. The reason is that the large prize justifies the high effort, and strategic complementarity reinforces this war effect. Large efforts nearly cancel out and the score is moved by randomness. Once a player attains a lead, he can first reinforce it and then maintain it using the threat that, should the opponent try to catch up, the game would enter again the “war” phase. A player can always react in time because the output difference evolves continuously. So both players maintain low efforts unless the score is almost a tie. This equilibrium pattern makes total expected output (or efforts) eventually decline with the prize, and is terrible for cost smoothing. The finite optimal prize makes for a more “interesting” race, fought all the way through.

Finally, we ask which is the optimal distortion of the score that the principal can introduce. Rather than deciding the game based on output difference, the principal can choose some non-linear function of this difference. This corresponds to the notion of point-by-point re-weighting of the score. We show how to solve in practice for the symmetric equilibrium strategy when players face any scoring function. Given the strategy, the principal’s objective is evaluated just as in the previous case. The choice of the optimal scoring rule appears to be non-recursive, and to not lend itself to any standard method of dynamic (infinite-dimensional) optimization. Computing a discrete approximation of any desired accuracy to the optimal scoring rule is a finite problem, but it appears that it can be solved only by trying all possibilities, namely, it is NP-complete.

We conjecture that the optimal scoring rule is concave: it is optimal to penalize the leader and help the follower as the game unfolds. Penalizing the leader, by making it harder and harder for him to move his score up, has undesirable ex ante incentives, because it reduces the returns to put in effort and gain a lead, but desirable ex post incentives, because it encourages the laggard to try harder and not to give up, as well as the leader to work really hard to preserve the acquired lead against the wind of the adverse scoring rule.

2 A Symmetric Linear-Quadratic Tug-of-War

2.1 Setup

Two agents $i = A, B$ compete for a prize $\pi > 0$. Agent i chooses continuously n_{it} at quadratic cost $c_0 n_{it}^2/2$, where $c_0 > 0$, to produce flow output

$$dx_{it} = n_{it}\mu dt + \frac{\sigma}{\sqrt{2}}dW_{it}$$

where the W_{it} are Wiener and mutually independent ($W_B \perp W_A$). Let $z_t \equiv x_{Bt} - x_{At}$, with

$$\begin{aligned} dz_t &= (n_{Bt} - n_{At})\mu dt + \frac{\sigma}{\sqrt{2}}d(W_{Bt} - W_{At}) \\ &= (n_{Bt} - n_{At})\mu dt + \sigma dW_t \end{aligned}$$

where W_t is Wiener and $z_0 = 0$. The rules of the competition are as follows: for some $\omega > 0$, player B wins if and when z_t hits $\omega > 0$, and player A wins if and when z_t hits $-\omega$. The winner gets $\pi > 0$ and the loser gets nothing. Payoffs are undiscounted.

3 Equilibrium Analysis

3.1 Equilibrium Definition

We restrict attention to Symmetric Markov Perfect Nash Equilibrium (SMPE), hereby defined on the one dimensional state space of score differences z .

Fix a pair of functions $n_i : [-\omega, \omega] \rightarrow R_+$, $i = A, B$ such that the following stochastic differential equation

$$dz_t = [n_A(z_t) - n_B(z_t)]\mu dt + \sigma dW_t$$

s.t. $z_0 = 0$ has a unique solution $z_t^{n_A, n_B, z_0}$. We define any such pair of functions *admissible*. En example is a pair of constant functions. Let

$$T_\omega(n_A, n_B, z_0) = \inf \{t \geq 0 : z_t^{n_A, n_B, z_0} = \omega\}$$

be the stopping time at which the process $z_t^{n_A, n_B, z_0}$ hits ω . A Markov Perfect Equilibrium is a pair of admissible functions $\{n_i(z)\}_{i=A,B}$ such that for all $z_0 \in (-\omega, \omega)$, $n_i(z_0)$ maximizes the expected prize minus costs given that at all other points player i will follow the strategy

n_i and the opponent $n_{-i}(z)$. For player B

$$n_B(z_0) = \arg \max_u \left\{ \pi \Pr(T_\omega(n_A, u, z_0) < T_{-\omega}(n_A, u, z_0)) \right. \\ \left. - E \left[\int_0^{T_{-\omega}(n_A, u, z_0) \vee T_\omega(n_A, u, z_0)} c_0 \frac{u_B^2(z_t^{n_A, u, z_0})}{2} dz_t^{n_A, u, z_0} \right] \right\}$$

and similarly for player A , where ω is replaced by $-\omega$ and A and B indexes are swapped.

We now impose symmetry, and look for a Symmetric MPE, a function n such that

$$n_B(z) = n_A(-z) \equiv n(z).$$

Assume that a SMPE exists and players play it. Let $\bar{V}_i(z)$ the continuation value to player i of the game starting from z in the SMPE path. By symmetry, we study the function

$$\bar{V}_B(z) = \bar{V}_A(-z) \equiv \bar{V}(z).$$

This $\bar{V}$ solves the HJB equation

$$0 = \sup_u -c_0 \frac{u^2}{2} + [u - n(-z)]\mu \bar{V}'(z) + \frac{\sigma^2}{2} \bar{V}''(z)$$

s.t. $\bar{V}(\omega) = \pi$ and $\bar{V}(-\omega) = 0$. In equilibrium, the maximizer u is $n(z)$, so the NFOC for player B is

$$u(z) = \frac{\mu}{c_0} \bar{V}'(z) \equiv n(z)$$

and for player A it is

$$n_A(z) = n(-z) = -\frac{\mu}{c_0} \bar{V}'_A(z) = \frac{\mu}{c_0} \bar{V}'(-z).$$

Plugging these NFOC back into the HJB equation yield the two-boundary value *Delayed Differential Equation* problem:

$$0 = \frac{[\mu \bar{V}'(z)]^2}{2c_0} - \frac{\mu^2}{c_0} \bar{V}'(-z) \bar{V}'(z) + \frac{\sigma^2}{2} \bar{V}''(z) \quad (1)$$

or

$$\bar{V}''(z) = \frac{\mu^2}{c_0 \sigma^2} \bar{V}'(z) [2\bar{V}'(-z) - \bar{V}'(z)] \equiv \delta \bar{V}'(z) [2\bar{V}'(-z) - \bar{V}'(z)] \quad (2)$$

s.t. $\bar{V}(-\omega) = 0$, $\bar{V}(\omega) = \pi$, hereby defining δ .

3.2 The Normalized Game

We can rescale the problem. Let

$$V(y) \equiv V(z/\omega) \equiv \delta \bar{V}(z)$$

hereby defining $y = z/\omega \in (-1, 1)$. Then (2) becomes

$$V''(y) = V'(y)[2V'(-y) - V'(y)]. \quad (3)$$

with boundary conditions

$$V(-1) = 0, V(1) = \delta\pi.$$

Therefore, WLOG, we can find the equilibrium value V for a normalized finish line $\omega = 1$, and then recover the value of the original problem at score z by

$$\bar{V}(z) = \frac{1}{\delta\pi} V\left(\frac{z}{\omega}\right)$$

and, from there, the SMPE strategies of the original game with finish line ω :

$$n_\omega(z) = \frac{\mu}{c_0} \bar{V}'(z) = \frac{\mu}{c_0 \delta\pi} \frac{1}{\omega} V'\left(\frac{z}{\omega}\right). \quad (4)$$

Notice that doubling ω and z halves effort n . That is, for every z and $y = z/\omega \in (-1, 1)$

$$\omega n_\omega(\omega y) = \frac{\mu}{c_0 \delta\pi} V'(y) = n_1(y)$$

so

$$n_\omega(z) = \frac{1}{\omega} n_1\left(\frac{z}{\omega}\right). \quad (5)$$

When the finish line doubles, players spend exactly half the effort at double the distance from the starting point $z_0 = 0$.

From now, WLOG we set $\omega = 1$ and solve the corresponding normalized problem where the only parameter is $\delta\pi$. Then we recover $n_\omega(z)$ from (5).

3.3 The Key Differential Equation

We can reduce the normalized problem to a *first-order* Delayed Differential Equation (DDE) in the effort best response, by letting $m(z) = V'(z)$, so that from (3)

$$m'(z) = m(z)[2m(-z) - m(z)] \quad (6)$$

subject to

$$\int_{-1}^1 m(z) dz = \delta\pi. \quad (7)$$

where we used $V(-1) = 0$ and $V(1) = \delta\pi$. We can transform the first order DDE into a:

Lemma 1 *The marginal value $m(z) = V'(z)$ obeys the second order ODE*

$$2m(z)m''(z) = 3(m'(z))^2 - 4m'(z)m(z)^2 - 3m(z)^4 \quad (8)$$

Proof: Take another derivative in (6)

$$m''(z) = 2 \{m'(z)[m(-z) - m(z)] - m(z)m'(-z)\}. \quad (9)$$

Equation (6) allows us to eliminate both $m(-z)$ and $m'(-z)$ from (9) using $2m(-z) = m'(z)/m(z) + m(z)$ and $m'(-z) = m(-z)[2m(z) - m(-z)]$, and deduce

$$\begin{aligned} m''(z) &= 2m'(z)[m(-z) - m(z)] - 2m(z)m'(-z) \\ &= 2m'(z) \left[m(-z) - \frac{1}{2}m(z) - \frac{1}{2}m(z) \right] - 2m(z)m(-z)[2m(z) - m(-z)] \\ &= 2m'(z) \left[\frac{m'(z)}{2m(z)} - \frac{1}{2}m(z) \right] - 2m(z) \left[\frac{m'(z)}{2m(z)} + \frac{1}{2}m(z) \right] \left[\frac{3}{2}m(z) - \frac{m'(z)}{2m(z)} \right] \\ &= \frac{3[m'(z)]^2 - 4m'(z)m^2(z) - 3m^4(z)}{2m(z)} \end{aligned}$$

Altogether this yields (8), as required. ■

If we can find a solution $m(z)$ to (8) subject to (7), we can recover the equilibrium strategy of the original game with $\omega = 1$ combining (4) and the definition $m = V'$:

$$n_1(z) = \frac{\sigma^2}{\mu} m(z) \quad (10)$$

From this function, we can recover the SMPE strategy of the original unscaled game with finish line ω

$$n_\omega(z) = \frac{\sigma^2}{\mu\omega} m\left(\frac{z}{\omega}\right).$$

In the Appendix we also derive a first-order ODE solved by the equilibrium strategy m , which has no explicit analytic solution but may be useful for some computations.

3.4 Characterization of the Equilibrium Strategy

We verify that our SMPE has the same qualitative properties as the equilibrium strategies of the Harris and Vickers (1987) tug-of-war, which is cast in discrete time and state space, with exponential chance of success in the stage game.

Lemma 2 *The SMPE strategy is strictly positive for all $z \in (-\omega, \omega)$:*

$$n_\omega(z) = \frac{\sigma^2}{\mu\omega} V' \left(\frac{z}{\omega} \right) > 0$$

Proof. The normalized HJB equation (3) reveals that if ever $V'(z) = m(z) = 0$, then we must have $V''(z) = m'(z) = 0$. But (3) can be differentiated as many times as desired, and each summand has a lower order derivative factor of V at each stage, so all derivatives of V vanish at z . For instance,

$$V'''(y) = V''(y)[2V'(-y) - V'(y)] + V'(y)[-2V''(-y) - V''(y)]. \quad (11)$$

But then V is constant in a neighborhood of y . Extending this, V is constant on the entire domain, contrary to the boundary conditions $V(-1) = 0 < \pi = V(1)$. ■

Lemma 3 (Rising Effort ... and Maybe Then Falling) *The SMPE strategy is either globally increasing or rising and then falling. In particular, it is increasing at $z = -\omega$.*

Proof. Equations (6) and (11) respectively yield

$$m'(0) = m^2(0) \quad \text{and} \quad m''(0) = -2m^3(0) < 0.$$

So at the beginning of the game, when it is tied at $z_0 = 0$, equilibrium effort is positive, and locally increasing and concave in the score. Since $m(z)$ is increasing at $z = 0$, it cannot be globally declining. Modifying (8), we find the equivalent second order nonlinear ODE

$$m''(z) = \frac{3}{2} \frac{(m'(z))^2}{m(z)} - 2m'(z)m(z) - \frac{3}{2}m(z)^3. \quad (12)$$

Assuming that $m'(z) = 0$, we find $m''(z) = -3m(z)^3/2 < 0$ because $m(z) > 0$. So either $m(z)$ is monotonic increasing or it has a unique maximum. ■

Lemma 4 (Follow the Leader) *The leader tries harder than the follower.*

Proof. Define $q(z) = m(z) - m(-z)$, with $q(0) = 0$. We claim that $q(z) \geq 0$ as $z \geq 0$.

$$\begin{aligned} q'(z) &= m'(z) + m'(-z) \\ &= m(z)[2m(-z) - m(z)] + m(-z)[2m(z) - m(-z)] \\ &= 4m(z)m(-z) - m^2(z) - m^2(-z) \end{aligned}$$

So $q'(0) = 2m^2(0) > 0$. Let z' be the least $z' \in (0, 1)$ with $q(z') = 0$. Then $m(z') = m(-z')$, and

$$q'(z') = 2m^2(z') > 0.$$

But then there is a smaller zero $z'' \in (0, z')$ of $q(z)$, contradicting the choice of z' . ■

In summary, in the Symmetric Markov Perfect Equilibrium,

- the strategy n_ω is “scalable” in the finish line ω according to (5), where n_1 is defined by (10) and m is the solution to the boundary value problem (7)–(12). the other parameters μ, σ, c_0, π matter only through the composite parameter $\delta\pi = \mu^2\pi / (c_0\sigma^2)$.
- the strategy n_ω is globally positive, and increasing or hump-shaped
- the leader always exerts more effort than the follower

4 The Mechanism Design Problem

4.1 Two Principal’s Objective Functions

Our main focus is on the optimal design of the race. Once we have characterized and computed the equilibrium strategy n , we can focus on the optimal design of the race. Let T denote the stopping time of the game when the finish line is ω and the prize is π , namely

$$T = \inf_{t \geq 0} \{z_t = \omega \text{ or } z_t = -\omega\}.$$

Denote the total expected output starting from z_0 :

$$\begin{aligned} Q_{\omega\pi}(z_0) &= E[x_{AT} + x_{BT} | z_0] \\ &= E \left[\int_0^T d(x_{At} + x_{Bt}) | z_0 \right] \\ &= E \left[\int_0^T (\mu(n_{At} + n_{Bt}) dt + \sigma dW_t) | z_0 \right] \end{aligned}$$

$$\begin{aligned}
&= \mu E \left[\int_0^T (n_{At} + n_{Bt}) dt | z_0 \right] \\
&= \frac{\mu}{\omega} E \left[\int_0^T \left(n_1 \left(\frac{z_t}{\omega} \right) + n_1 \left(-\frac{z_t}{\omega} \right) \right) dt | z_0 \right]
\end{aligned}$$

by the effort scaling relationship (5), and the martingale property of the Wiener noise. Similarly, let

$$K_{\omega, \pi}(z_0) = \frac{c_0}{2} E \left[\int_0^T (n_{At}^2 + n_{Bt}^2) dt | z_0 \right] = \frac{c_0}{2\omega} E \left[\int_0^T \left(n_1^2 \left(\frac{z_t}{\omega} \right) + n_1^2 \left(-\frac{z_t}{\omega} \right) \right) dt | z_0 \right]$$

denote the expected total costs of effort during the game.

The problem of the principal can be taken to be $\max_{\pi, \omega} \{Q_\omega(0) - \pi\}$ or $\max_{\pi, \omega} \{U_\omega(0) - \pi\}$. The optimal (effort) value Q solves $Q_\omega(-\omega) = Q_\omega(\omega) = 0$, and the HJB equation

$$0 = n_\omega(z) + n_\omega(-z) + [n_\omega(z) - n_\omega(-z)] Q'_\omega(z) + \frac{\sigma^2}{2\mu} Q''_\omega(z) \quad (13)$$

Similarly, expected costs K solve $K_\omega(-\omega) = K_\omega(\omega) = 0$ and

$$0 = \frac{c_0}{2} \{n_\omega^2(z) + n_\omega^2(-z)\} + [n_\omega(z) - n_\omega(-z)] K'_\omega(z) \mu + \frac{\sigma^2}{2} K''_\omega(z)$$

Finally, net output $U_\omega = Q_\omega - K_\omega$ solves $U_\omega(-\omega) = U_\omega(\omega) = 0$ and

$$0 = -\frac{c_0}{2} \{n_\omega^2(z) + n_\omega^2(-z)\} + \{n_\omega(z) + n_\omega(-z) + [n_\omega(z) - n_\omega(-z)] U'_\omega(z)\} \mu + \frac{\sigma^2}{2} U''_\omega(z)$$

Given the equilibrium strategy n_ω , these functions solve linear ordinary differential equations. so we can find them using the standard method of variation of the arbitrary constant and the fact that Q and U are even functions, and thus flat at the origin. Omitting the tedious derivation, we obtain:

$$Q_\omega(z) = \frac{2\mu}{\sigma^2} \int_{-\omega}^z \left[\int_j^0 \frac{n_\omega(y) + n_\omega(-y)}{\exp \left\{ \frac{2\mu}{\sigma^2} \int_j^y [n_\omega(-x) - n_\omega(x)] dx \right\}} dy \right] dj$$

and

$$U_\omega(z) = \frac{2\mu}{\sigma^2} \int_{-\omega}^z \int_j^0 \frac{[n_\omega(y) + n_\omega(-y)] - \frac{c_0}{2\mu} [n_\omega^2(y) + n_\omega^2(-y)]}{\exp \left\{ \frac{2\mu}{\sigma^2} \int_j^y [n_\omega(-x) - n_\omega(x)] dx \right\}} dy dj.$$

4.2 Homogeneity of the Principal's Value

One can immediately see that a homogeneity of the principal's value function in the finish line ω . Fix ω . For $z \in [-\omega, \omega]$, we guess

$$Q_\omega(z) \equiv \omega Q_1\left(\frac{z}{\omega}\right)$$

where, by definition, $Q_1(z/\omega)$ solves

$$0 = \left\{ n_1\left(\frac{z}{\omega}\right) + n_1\left(-\frac{z}{\omega}\right) + \left[n_1\left(\frac{z}{\omega}\right) - n_1\left(-\frac{z}{\omega}\right) \right] Q_1'\left(\frac{z}{\omega}\right) \right\} \mu + \frac{\sigma^2}{2} Q_1''\left(\frac{z}{\omega}\right)$$

subject to $Q_1(-1) = Q_1(1) = 0$. Substituting $Q'_\omega(z) = Q'_1(z/\omega)$, $Q''_\omega(z) = Q''_1(z/\omega)/\omega$ and $n_1\left(\frac{z}{\omega}\right) = \omega n_\omega(z)$ from (5) in the last ODE for $z \in [-\omega, \omega]$ yields

$$\begin{aligned} 0 &= \left\{ \omega n_\omega(z) + \omega n_\omega(-z) + [\omega n_\omega(z) - \omega n_\omega(-z)] Q'_\omega(z) \right\} \mu + \frac{\sigma^2}{2} \omega Q''_\omega(z) \\ &= \left\{ n_\omega(z) + n_\omega(-z) + [n_\omega(z) - n_\omega(-z)] Q'_\omega(z) \right\} \mu + \frac{\sigma^2}{2} Q''_\omega(z) \end{aligned}$$

so $Q_\omega(z)$ solves (13) for $n(z) = n_\omega(z)$ and $z \in [-\omega, \omega]$. In addition

$$\begin{aligned} Q_\omega(-\omega) &= -\omega Q_1\left(-\frac{\omega}{\omega}\right) = -\omega Q_1(-1) = -\omega \cdot 0 = 0 \\ Q_\omega(\omega) &= \omega Q_1\left(\frac{\omega}{\omega}\right) = \omega Q_1(1) = \omega \cdot 0 = 0 \end{aligned}$$

so Q_ω also solves the appropriate boundary conditions at $z = \pm\omega$. Therefore the guess is verified.

In particular

$$Q_\omega(0) = \omega Q_1(0).$$

Since $Q_1(0) > 0$, expected output at the beginning of the game is linear and unbounded in the finish line ω . The finish line ω has a total effect on the principal's value, combining a direct one and an indirect one through equilibrium strategies, which is overall linear. So, there exists no optimal finish line. Since this appears to be an artifice of the linear-quadratic setup, which has otherwise many advantages, we set $\omega = 1$ and we restrict attention to the choice of the optimal prize π .

For effort costs, we guess

$$K_1\left(\frac{z}{\omega}\right) = K_\omega(z)$$

plugging back and using again $n_1\left(\frac{z}{\omega}\right) = \omega n_\omega(z)$

$$0 = \frac{c_0}{2\omega^2} \left\{ n_1^2\left(\frac{z}{\omega}\right) + n_1^2\left(-\frac{z}{\omega}\right) \right\} + \frac{1}{\omega} \left[n_1\left(\frac{z}{\omega}\right) - n_1\left(-\frac{z}{\omega}\right) \right] \frac{1}{\omega} K_1'\left(\frac{z}{\omega}\right) \mu + \frac{\sigma^2}{2\omega^2} K_\omega''(z)$$

simplifying ω^2 we see that the guess is verified, and the boundary conditions hold. So

$$K_\omega(0) = K_1(0)$$

the expected effort costs are independent of the finish line: players spread themselves thinner on a wider game domain. Clearly, if the expected squared effort K_ω is constant in ω the expected effort Q_ω must be increasing. So

$$U_\omega(0) = \omega Q_1(0) - K_1(0)$$

output net of effort costs is linear in the finish line.

4.3 The Optimal Prize

We look for a solution m to the parameter-free ODE (12) subject to

$$\int_{-1}^1 m(z) dz = \delta\pi.$$

We then plug this function into the principal's objective at $z = 0$ and we use

$$\begin{aligned} m(z) &= \frac{\mu}{\sigma^2} n(z) \\ \frac{c_0}{\sigma^2} \left(\frac{\sigma}{\mu} \right)^2 &= \frac{1}{\delta} \end{aligned}$$

to write ex ante expected output gross, as follows:

$$Q(0) = 2 \int_{-1}^0 \left[\int_j^0 \frac{m(y) + m(-y)}{\exp \left\{ 2 \int_j^y [m(-x) - m(x)] dx \right\}} dy \right] dj \quad (14)$$

$$U(0) = \int_{-1}^0 \int_j^0 \frac{2[m(y) + m(-y)] - \delta^{-1}[m^2(y) + m^2(-y)]}{\exp \left\{ 2 \int_j^y [m(-x) - m(x)] dx \right\}} dy dj. \quad (15)$$

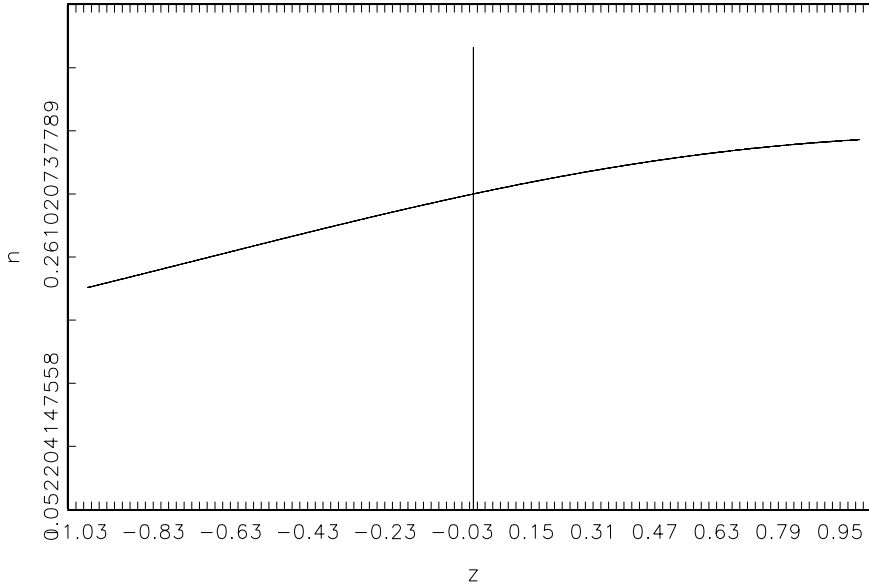
Since the equilibrium strategy m only depends on parameters through $\delta\pi$, we have reduced the parameters of the problem from four (μ, σ, c_0, π) to two: $\delta = \frac{\mu^2}{\sigma^2 c_0}$ and π .

Conjecture 5 *The prize π^* that maximizes expected efforts gross of costs, $Q(0)$, is finite.*

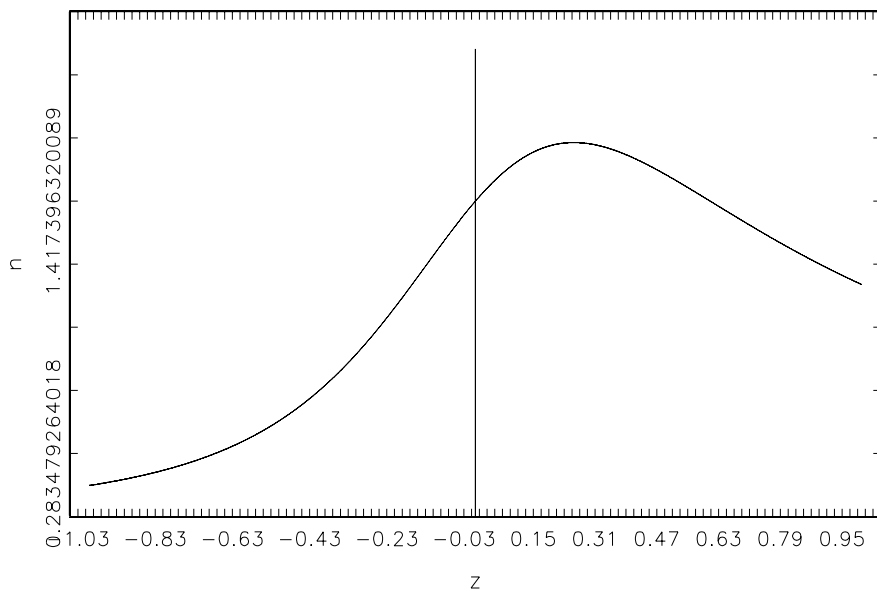
This result — verified in some numerical examples — stands in stark contrast with the standard war of attrition, in which all rents are competed away. In other words, a higher prize yields greater effort. Here, the equilibrium strategy is unimodal and converges to a spike at $z = 0$ of increasing height. Players engage in an initial “war” phase, where their efforts are very large near $z = 0$. Eventually, we know that efforts eventually taper off.

The figures show the equilibrium strategy $n(z)$ computed for a particular parameter configuration. We can see that, as proven in general, the strategy is either increasing or, when the prize is large enough, hump-shaped. At the optimal prize which maximizes the principal’s objective, initially expected output $Q_1(0)$, the players’ equilibrium strategy $n(z)$ is hump-shaped.

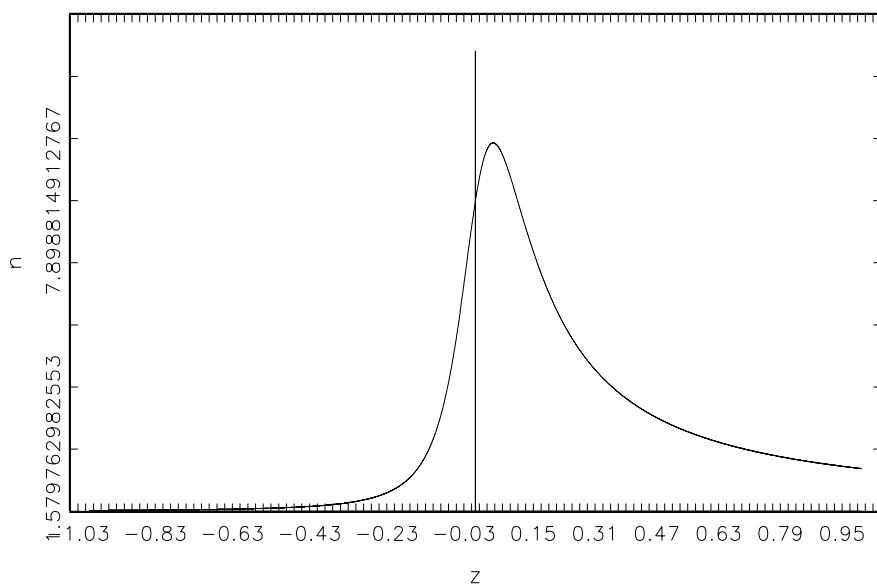
$\mu/\sigma=2.0$ cost scale=1.0 and SUBoptimal prize is 1.02000



mu/sigma=2.0 cost scale=1.0 and optimal prize is 4.02000



mu/sigma=2.0 cost scale=1.0 and (too large) prize is 8.02000



Surprisingly, the optimal prize is finite even if it has no opportunity cost. Contrast this result to a one-agent incentive problem, where the agent is motivated to exert effort with convex cost by either a flow payment or a prize once he achieves a certain target. Clearly,

as long as the principal cares only about expected output, and neither about the agent's effort cost nor the opportunity cost of the prize, the optimal payment to the single agent is unbounded above. In the race, conversely, too large a prize is detrimental to the principal.

The numerical examples, in particular the third figure, suggest the reason for and an interpretation of this result. When the prize rises beyond a point, the equilibrium strategy is unimodal and converges to a spike near $z = 0$, of increasing height. Players engage in an initial “war” phase, where their efforts are very large near $z = 0$. The reason is that the prize justifies the high effort, and strategic complementarity reinforces that. Large efforts nearly cancel out and the score is moved by randomness. Once a player attains a lead, he can first reinforce it and then maintain it using the threat that, should the opponent try to catch up, the game would enter again the “war” phase. A player can always react in time because the score z evolves continuously. So both players maintain low efforts unless the score is almost a tie. Notice from the third figure that the difference in efforts $n(z) - n(-z)$, which determines the mean speed of the score z , rises extremely fast with z initially. So the war phase is very short-lived, and most of race is fought at relatively low intensity. This equilibrium pattern makes output rise modestly with the prize, and is terrible for cost smoothing. The finite optimal prize makes for a more “interesting” race, fought all the way through.

Comparing the three figures, one also notices the rising scale of effort with the prize. One may think that higher effort is better for the principal. But too high effort ends the game too quickly. The principal does not care about the area under the effort strategy in the figures, but rather about the expected effort during the entire duration of the game, $Q(0)$. The expression (14) for the principal's objective $Q(0)$ makes this point transparent: before integrating over the game domain, the sum of efforts is divided by an exponential function of the difference in efforts, $\exp \left\{ 2 \int_j^y [m(-x) - m(x)] dx \right\}$. Since we are integrating over $x < 0$ to exploit symmetry, the difference $m(-x) - m(x)$ is that between the leader's and the laggard's efforts, that we have proven in general to be positive. By assumption, this difference determines the speed of the score z and thus the duration of the game. The higher the leader's effort relative to the laggard's, the larger is this exponential term in the denominator of $Q(0)$, the stronger is this effect detrimental to the principal. In the third figure, it is clear that this difference in efforts is indeed very large over much of the game domain.

For comparative statics, we conjecture that weaker players, either because of higher effort costs $\bar{c}_0$ or lower productivity of effort μ/σ , must be motivated by a higher optimal prize. Since the parameter space has been reduced to δ and π , we will be able to verify this conjecture numerically WLOG. We also plan to establish analytic results.

5 The Optimal Scoring Function

The equilibrium of the game often features a declining effort of the leader when the score rises above some threshold. In fact, the equilibrium strategy is either globally increasing or hump-shaped, and the latter case arises when the normalized prize $\delta\pi$ is large enough, including when it is chosen optimally. As the lead widens further, the leader's effort declines further, and the laggard's effort also declines, so the sum of the two efforts declines. This outcome runs counter the principal's objective. The strategic force behind this undesirable outcome is simple. Once the lead is large, the laggard is discouraged, because catching up would require a large amount of concentrated effort, that would also be matched by the leader ramping up his defense. As the prospects of victory become dim, the laggard's effort declines, and the leader can afford slacking off, under the threat of responding harshly to any attempt by the laggard to catch up.

In order to ameliorate this problem, the principal should prefer to distort incentives directly during the game, not only indirectly through the choice of the prize at the endpoint. Intuitively, it may be optimal to penalize the leader by making his effort less productive, to keep the laggard hopeful and in the race. The trade-off is clear, as penalizing the leader has also detrimental effects on the incentives to fight and gain a lead in the first place.

We now address the more general mechanism design question of how incentives should be set directly during the game, by re-weighting the performance measure z into a score s . We show how to set up and solve the problem of an optimal scoring function.

5.1 Equilibrium of the Race For Given Scoring Function

Suppose now that the principal designs the rules of the game in terms of a "score"

$$s = f(z) \tag{16}$$

so that player B wins if $s = \omega$. We aim to solve for the SMPE of this game n^f . Then the principal solves the same problem as before, namely, he cares about actual output and effort costs, not about the score per se, so he maximizes either $Q(0)$ or $U(0)$, where the strategy n^f replaces n . Since the principal can always set $s = z$, as in the previous case, and obtain an objective that grows linearly in ω , to make the problem nontrivial we still set $\omega = 1$. For simplicity, we also omit the superscript f and let $n : [-1, 1] \rightarrow R_+$ be the equilibrium strategy in score space.

We assume that f is a C^2 function, so that by Ito's lemma

$$ds_t = f'(z_t) (n_{At} - n_{Bt}) \mu dt + \sigma f'(z_t) dW_t + \frac{\sigma^2 f''(z_t)}{2} dt.$$

Conjecture 6 (Monotone Incentives) *The optimal scoring function f is increasing.*

If it was not, then player A would like to let output decline locally to give the score the best chance to rise, by setting effort to zero (effort cannot be negative, one cannot “intentionally score an own goal,” say.) Zero effort is detrimental to the principal.

If f is increasing and thus invertible, then

$$\phi(s) \equiv f'(f^{-1}(s)) = f'(z) \tag{17}$$

so that

$$\phi'(s) = \frac{f''(f^{-1}(s))}{f'(f^{-1}(s))} = \frac{f''(z)}{\phi(s)}$$

and replacing

$$ds_t = \phi(s_t) \left\{ \left[(n_{At} - n_{Bt}) \mu + \frac{\sigma^2}{2} \phi'(s_t) \right] dt + \sigma dW_t \right\}$$

The winner is the first to hit score $s = 1$ or $s = -1$.

The HJB equation in a SMPE of this distorted game is

$$0 = \max_u -c_0 \frac{u^2}{2} + \mu [u - n(-s)] \phi(s) \bar{V}'(s) + \frac{\sigma^2 \phi'(s) \phi(s)}{2} \bar{V}'(s) + \frac{\sigma^2}{2} \phi^2(s) \bar{V}''(s).$$

The NFOC for optimal effort is

$$c_0 n(s) = \mu \phi(s) \bar{V}'(s)$$

substituting back into the HJB equation, rearranging and using the definition of δ

$$\begin{aligned}
\bar{V}''(s) &= \frac{\mu^2}{c_0 \sigma^2} \left[2 \frac{\phi(-s)}{\phi(s)} \bar{V}'(-s) - \bar{V}'(s) \right] \bar{V}'(s) - \frac{\phi'(s)}{\phi(s)} \bar{V}'(s) \\
&= \delta \left[2 \frac{\phi(-s)}{\phi(s)} \bar{V}'(-s) - \bar{V}'(s) \right] \bar{V}'(s) - \frac{\phi'(s)}{\phi(s)} \bar{V}'(s).
\end{aligned}$$

If $s = z$, we have $\phi(s) = 1$ and $\phi'(s) = 0$, so this reduces to the familiar equation (2).

Once again, normalize the value function

$$V(z) = \delta \bar{V}(z)$$

so that the value of the game to the player solves

$$\begin{aligned}
\frac{V''(s)}{\delta} &= \delta \left[2 \frac{\phi(-s)}{\phi(s)} \frac{V'(-s)}{\delta} - \frac{V'(s)}{\delta} \right] \frac{V'(s)}{\delta} - \frac{\phi'(s)}{\phi(s)} \frac{V'(s)}{\delta} \\
V''(s) &= \left[2 \frac{\phi(-s)}{\phi(s)} V'(-s) - V'(s) \right] V'(s) - \frac{\phi'(s)}{\phi(s)} V'(s)
\end{aligned}$$

subject to

$$V(-1) = 0 \text{ and } V(1) = \delta \pi.$$

In terms of strategies, let

$$m(s) = \phi(s) V'(s)$$

which clearly depends only on the scoring function through ϕ and on $\delta \pi$ through V . As before, we will derive an ODE for m . As in the original, undistorted game, from the solution $m(s)$ we can then recover the equilibrium strategy of the original game:

$$n(s) = \frac{\mu}{c_0} \phi(s) \bar{V}'(s) = \frac{\mu}{c_0} \phi(s) \frac{V'(s)}{\delta} = \frac{\mu}{c_0 \delta} \phi(s) V'(s) = \frac{\mu}{c_0 \delta} m(s).$$

Plugging back this effort strategy into the principal's objective, we see that the latter depends only on the composite parameter δ and on the prize π , for every scoring function ϕ .

Differentiating m

$$m'(s) = \phi'(s) V'(s) + \phi(s) V''(s)$$

we get

$$\begin{aligned}
V'(s) &= \frac{m(s)}{\phi(s)} \\
V''(s) &= \frac{m'(s)}{\phi(s)} - \frac{\phi'(s)}{\phi(s)} V'(s) = \frac{m'(s)}{\phi(s)} - \frac{\phi'(s)}{\phi(s)} \frac{m(s)}{\phi(s)}
\end{aligned}$$

plugging back into the HJB equation

$$\frac{m'(s)}{\phi(s)} - \frac{\phi'(s)}{\phi(s)} \frac{m(s)}{\phi(s)} = \left[2 \frac{\phi(-s)}{\phi(s)} \frac{m(-s)}{\phi(-s)} - \frac{m(s)}{\phi(s)} - \frac{\phi'(s)}{\phi(s)} \right] \frac{m(s)}{\phi(s)}.$$

After simplifying, we obtain a DDE in effort m ,

$$m'(s)\phi(s) = 2m(s) \left[m(-s) - \frac{1}{2}m(s) \right] \quad (18)$$

which is identical to (6), except for the $\phi(s)$ on the LHS.

To turn (18) into an Ordinary DE, we use its two equivalent forms.

$$\begin{aligned} m(-s) &= \frac{1}{2}m(s) + \frac{m'(s)\phi(s)}{2m(s)} \\ m'(-s) &= 2 \frac{m(-s)}{\phi(-s)} \left[m(s) - \frac{1}{2}m(-s) \right] \end{aligned}$$

We proceed as in Lemma 1: differentiating once more (18)

$$m''(s)\phi(s) + m'(s)\phi'(s) = 2 \{ m'(s)[m(-s) - m(s)] - m(s)m'(-s) \}.$$

Substituting from the DDE we get

$$\begin{aligned} \phi(s)m''(s) + m'(s)\phi'(s) &= 2m'(s) \left[m(-s) - \frac{1}{2}m(s) - \frac{1}{2}m(s) \right] - 2m(s) \cdot 2 \frac{m(-s)}{\phi(-s)} \left[m(s) - \frac{1}{2}m(-s) \right] \\ &= 2m'(s) \left[\frac{m'(s)\phi(s)}{2m(s)} - \frac{1}{2}m(s) \right] - 2m(s) \cdot 2 \frac{\frac{1}{2}m(s) + \frac{m'(s)\phi(s)}{2m(s)}}{\phi(-s)} \left\{ m(s) - \frac{1}{2} \left[\frac{1}{2}m(s) + \frac{m'(s)\phi(s)}{2m(s)} \right] \right\} \\ &= \frac{[m'(s)]^2 \phi(s)}{m(s)} - m'(s)m(s) - m(s) \frac{m(s) + \frac{m'(s)\phi(s)}{m(s)}}{\phi(-s)} \left\{ 2m(s) - \frac{1}{2}m(s) - \frac{m'(s)\phi(s)}{2m(s)} \right\} \\ &= \frac{[m'(s)]^2 \phi(s)}{m(s)} - m'(s)m(s) - \frac{m(s)}{\phi(-s)} \left\{ m(s) + \frac{m'(s)\phi(s)}{m(s)} \right\} \left\{ \frac{3}{2}m(s) - \frac{m'(s)\phi(s)}{2m(s)} \right\} \\ &= \frac{[m'(s)]^2 \phi(s)}{m(s)} - m'(s)m(s) - \frac{m(s)}{\phi(-s)} \left\{ \frac{3}{2}m^2(s) + \frac{3}{2}m'(s)\phi(s) - \frac{m'(s)\phi(s)}{2} - \frac{1}{2} \left[\frac{m'(s)\phi(s)}{m(s)} \right]^2 \right\} \\ &= \frac{[m'(s)]^2 \phi(s)}{m(s)} - m'(s)m(s) - \frac{m(s)}{\phi(-s)} \left\{ \frac{3}{2}m^2(s) + m'(s)\phi(s) - \frac{1}{2} \left[\frac{m'(s)\phi(s)}{m(s)} \right]^2 \right\} \\ &= \frac{[m'(s)]^2 \phi(s)}{m(s)} - m'(s)m(s) - \frac{3}{2} \frac{m^3(s)}{\phi(-s)} - m(s)m'(s) \frac{\phi(s)}{\phi(-s)} + \frac{1}{2} \frac{[m'(s)\phi(s)]^2}{m(s)\phi(-s)} \end{aligned}$$

Finally, the normalized equilibrium strategy of the distorted game is:

$$\phi(s)m''(s) = -m'(s)\phi'(s) + \frac{[m'(s)]^2 \phi(s)}{m(s)} \left[1 + \frac{\phi(s)}{2\phi(-s)} \right] - m'(s)m(s) \left[1 + \frac{\phi(s)}{\phi(-s)} \right] - \frac{3}{2} \frac{m^3(s)}{\phi(-s)} \quad (19)$$

subject to

$$\int_{-1}^1 V'(s)ds = \int_{-1}^1 \frac{m(s)}{\phi(s)}ds = \delta\pi. \quad (20)$$

In the natural score scale $s = z$ we have $\phi(s) = 1$, $\phi'(s) = 0$ and this boundary value problem reduces to (12).

With an affine scoring function $f(z) = a + bz$ we have $\phi(s) = f'(z) = b$ and $\phi'(s) = 0$. Let

$$u(s) = \frac{m(s)}{b}$$

then replacing into (19)-(20)

$$\begin{aligned} bu''(s) &= \frac{[bu'(s)]^2 b}{bu(s)} \left[1 + \frac{b}{2b} \right] - bu'(s)bu(s) \left[1 + \frac{b}{b} \right] - \frac{3}{2} \frac{b^3 u^3(s)}{b} \\ \frac{u''(s)}{b} &= \frac{3}{2} \frac{[u'(s)]^2}{u(s)} - 2u'(s)u(s) - \frac{3}{2} u^3(s) \end{aligned}$$

s.t.

$$\int_{-1}^1 u(s)ds = \delta\pi.$$

So the score location a is obviously irrelevant, but the score scale b is relevant.

Given all parameters, as summarized by δ , a scoring function f and the implied marginal score function $\phi(s) \equiv f'(f^{-1}(s))$, the SMPE of the contest is the function m that solves (19) subject to (20). Given this m , the principal obtains payoffs Q or U as defined in (14) or (15). The principal chooses the optimal prize π and scoring function f to maximize his payoff.

Conjecture 7 *The optimal scoring function f is concave.*

This conjecture says that it is optimal to penalize the leader, as the marginal net output z that he produces translates into lower and lower marginal score.

5.2 Computation of the Optimal Scoring Function

Clearly, the problem of finding the optimal f is non-recursive.

For any given scoring function f we can solve (19)-(20) equation numerically, by computing the numerical equivalent of ϕ and ϕ' , and using as “pivot”

$$m'(0) = \frac{m^2(0)}{\phi(0)}.$$

We guess $m(0)$, compute $m'(0)$, then $m''(0)$ from the ODE above, and extend the solution to $[-1, 1]$. Then we check whether the boundary condition (20) is satisfied, otherwise we iterate. A brute force algorithms considers and evaluated for the principal all increasing functions onto $[-1, 1]$. If f is discretized into a vector of size N , there is a finite set of such functions to be checked.

Example 1 *Power Scoring Functions.* *A restricted class of scoring functions is*

$$s = z^\alpha$$

for some α odd, so this power function is monotone and thus invertible. Then

$$\begin{aligned} f^{-1}(s) &= s^{\frac{1}{\alpha}} \\ \phi(s) &= \alpha s^{\frac{\alpha-1}{\alpha}} \\ \phi'(s) &= (\alpha - 1) s^{-\frac{1}{\alpha}} \end{aligned}$$

We use these expressions and search for the optimal α . In particular, we can ask is the optimal α positive or negative? I.e., is the leader encouraged or discouraged? Here we only need to look over a relatively small number of possible values of α .

References

Harris, Christopher and John Vickers. 1987. “Racing with Uncertainty.” *Review of Economic Studies*, 54(1), pp. 1-21

A Appendix. A 1st order ODE for the Equilibrium Strategy

Eqs. (6) and (12) are reported here.

$$\begin{aligned} m'(z) &= m(z)[2m(-z) - m(z)] \\ 2m(z)m''(z) &= 3(m'(z))^2 - 4m'(z)m^2(z) - 3m^4(z) \end{aligned}$$

The latter has an exact implicit solution

$$z = \int_0^{m(z)} h(s, C_1) ds - C_2$$

where

$$h(s, C_1) = s^{-2} \left[\sqrt[3]{-\frac{C_1}{s} + \frac{1}{9} \frac{C_1}{s} \sqrt{\left(-\frac{3}{s} C_1 + 81\right)}} + \frac{\frac{C_1}{s}}{3 \sqrt[3]{-\frac{C_1}{s} + \frac{1}{9} \frac{C_1}{s} \sqrt{\left(-\frac{3}{s} C_1 + 81\right)}}} - 3 \right]^{-1}.$$

Let

$$m_0 \equiv m(0).$$

Then

$$0 = \int_0^{m(0)} h(s, C_1) ds - C_2.$$

Solving for C_2 and substituting

$$z = \int_{m(0)}^{m(z)} h(s, C_1) ds.$$

Differentiating this implicit solution once yields

$$\begin{aligned} 1 &= m'(z) h(m(z), C_1) \\ h(m(z), C_1) &= \frac{1}{m'(z)} \end{aligned}$$

Using the DDE at $z = 0$

$$m'(0) = m^2(0) = m_0^2$$

so

$$h(m_0, C_1) = m_0^{-2}$$

Plugging back in the definition of h and solving for C_1 as a function of m_0 yields

$$C_1 = 32m_0$$

so we can replace $C_1 = 32m_0$ to get the implicit solution without constants of integration:

$$z = \int_{m_0}^{m(z)} s^{-2} \left[2 \sqrt[3]{-\frac{4m_0}{s} + \frac{4m_0}{s} \sqrt{1 - \frac{32}{81} \frac{m_0}{s}}} - 3 + \frac{16 \frac{m_0}{s}}{3 \sqrt[3]{-\frac{4m_0}{s} + \frac{4m_0}{s} \sqrt{1 - \frac{32}{81} \frac{m_0}{s}}}} \right]^{-1} ds. \quad (21)$$

Change variable of integration

$$x \equiv \frac{4m_0}{s}$$

so

$$dx = -\frac{4m_0}{s^2}ds$$

$$z = \frac{1}{4m_0} \int_{4\frac{m_0}{m(z)}}^4 \left[2\sqrt[3]{-x + x\sqrt{1 - \frac{8}{81}x}} - 3 + \frac{4x}{3\sqrt[3]{-x + x\sqrt{1 - \frac{8}{81}x}}} \right]^{-1} dx$$

Let

$$g(x) \equiv \int \left[2\sqrt[3]{-x + x\sqrt{1 - \frac{8}{81}x}} - 3 + \frac{4x}{3\sqrt[3]{-x + x\sqrt{1 - \frac{8}{81}x}}} \right]^{-1} dx$$

so that

$$4m_0z = g(4) - g\left(\frac{4m_0}{m(z)}\right)$$

and if g is invertible

$$m(z) = \frac{4m_0}{g^{-1}(4m_0z - g(4))}.$$

Consider again (21). Differentiate on both sides and rearranging yields the desired first-order ODE in the strategy m on $[-1, 1]$:

$$\frac{m'(z)}{m^2(z)} = 2\sqrt[3]{-\frac{4m_0}{m(z)} + \frac{4m_0}{m(z)}\sqrt{1 - \frac{32}{81}\frac{m_0}{m(z)}}} - 3 + \frac{16\frac{m_0}{m(z)}}{3\sqrt[3]{-\frac{4m_0}{m(z)} + \frac{4m_0}{m(z)}\sqrt{1 - \frac{32}{81}\frac{m_0}{m(z)}}}}. \quad (22)$$

subject to the isoperimetric constraint

$$\int_{-1}^1 m(z)dz = \delta\pi$$

which pins down m_0 . (22) can be simplified further. Let

$$p(z) \equiv -\frac{4m_0}{m(z)}$$

$$p'(z) = 4\frac{m_0m'(z)}{m^2(z)}$$

so (22) reads

$$\frac{p'(z)}{4m_0} = 2\sqrt[3]{p(z) - p(z)\sqrt{1 + \frac{8}{81}p(z)}} - \frac{4p(z)}{3\sqrt[3]{p(z) - p(z)\sqrt{1 + \frac{8}{81}p(z)}}} - 3. \quad (23)$$

The same result can be obtained with a “depressed cubic” approach.