

Uncertainty-driven cooperation

DORUK CETEMEN
Collegio Carlo Alberto

ILWOO HWANG
Department of Economics, University of Miami

AYÇA KAYA
Department of Economics, University of Miami

We consider dynamic team production in the presence of uncertainty. Team members receive interim feedback that depends on both their current effort level and the project's uncertain prospects. In this environment, each member can encourage the others by making them more optimistic about the project's prospects. We study the extent to which this incentive counters the usual free-riding incentive. Restricting the agents' access to feedback can increase their equilibrium effort levels by mitigating the ratchet effect. In this case, using joint performance measures can be beneficial even when individual measures are available.

KEYWORDS. Team production, free-riding, uncertainty, learning.

JEL CLASSIFICATION. C72, C73, D23, D83.

1. INTRODUCTION

Teams, as agile and adaptable units of production, are commonly utilized in the modern workplace. Businesses are increasingly moving away from traditional hierarchical structures to networks of project-based teams (Deloitte 2016). According to Lazear and Shaw (2007), the percentage of large firms utilizing self-managed teams in the United States is close to 80%. Similarly, Bandiera et al. (2013) cite evidence that 47% of British establishments have more than 90% of their workers organized in teams.¹

What distinguishes a team from a mere group of workers is “shared responsibility of work outcomes” (Hackman 1987). The economics literature has long recognized that

Doruk Cetemen: doruk.cetemen@carloalberto.org

Ilwoo Hwang: ihwang@bus.miami.edu

Ayça Kaya: akaya@bus.miami.edu

We are grateful to the anonymous referees for their insightful suggestions. We thank Paulo Barelli, Dan Bernhardt, Simon Board, Ralph Boleslavsky, Alessandro Bonatti, Yeon-Koo Che, Gonzalo Cisternas, George Georgiadis, Ben Golub, Hari Govindan, Yingni Guo, Marina Halac, Navin Kartik, George Mailath, Dmitry Orlov, Romans Pancs, Heikki Rantakari, Vasiliki Skreta, Takuo Sugaya, Curtis Taylor, Yuichi Yamamoto, Huseyin Yildirim, and various seminars and conferences audiences for insightful comments. This paper is a revised version of Chapter 1 of Cetemen's Ph.D. dissertation.

¹The cited source is “Workplace Employment Relations Survey” (2004). See footnote 1 in Bandiera et al. (2013) for a description of the source.

such arrangements are problematic due to inherent free-riding incentives (Hölmstrom 1982). On the positive side, with shared responsibility, each member of a team benefits from the hard work of others and, therefore, would like to encourage his teammates to contribute. This feature may be an advantage of using teams and potentially a reason why businesses use teams so often. When the environment affords tools and channels by which the team members can indeed encourage each other, the cost of free-riding can be more than justified by the additional value thereby created. In this paper, we show that such a counter force exists if the team interacts over time and the environment features uncertainty.² Moreover, we show how the organization can further exploit the encouragement mechanism by designing performance measures and feedback rules.

We consider a team of agents working over a finite horizon, on a joint project whose true prospects are unknown. Each agent's effort level is unobservable by the others. Over time, the agents receive interim public feedback about team performance; then they can adapt to new information by adjusting their effort levels. This feedback is noisy but informative about the agents' efforts and the project's potential: both high effort and good prospects statistically improve feedback. When the project ends, agents share the total output in a prespecified way.

Our main finding is a mutual encouragement effect among team members, associated with learning about the project's potential. In environments with uncertainty, team members' optimism plays an important role in how hard they work. Optimistic agents exert more effort as they expect higher returns for it. At the same time, interim feedback about performance affects the members' optimism. Given that a team member's effort can affect interim feedback, each member has incentives to work harder to preclude setbacks and, thus, keep her *teammates* optimistic, encouraging them to exert effort. By doing this, she encourages higher effort on *their* part. The possibility of such mutual encouragement, which we call *encouragement via belief manipulation*, leads to an increased equilibrium effort level that counters the free-riding problem. Importantly, the encouragement effect in our model is present only when there is uncertainty about the prospects of the project. This suggests that introducing uncertainty into team production can be welfare improving.³

Similar insights regarding the role of uncertainty in incentivizing agents to take desirable actions have appeared elsewhere, notably in the career concerns literature pioneered by Hölmstrom (1999). Ours is a novel application of such an encouragement mechanism in a team context, which is further distinguished because it features *mutual* (as opposed to one-sided) encouragement among team members.

Our analysis not only demonstrates that encouragement via belief manipulation can counter free-riding incentives, but also identifies a force that limits its benefits. When an agent engages in this type of encouragement, she creates optimism among her teammates, which in turn translates into high expectations of *her own* future effort. However, the agent does not share her teammates' optimism and, therefore, chooses not to meet

²With few notable exceptions, the existing literature considers either static environments or dynamic environments with complete information, which explains why this channel has not previously appeared.

³This result is related to Hermalin (1998), which shows that providing only one member (leader) superior information could lead to a better outcome compared with complete information.

their expectations. Therefore, optimism generated via belief manipulation is inherently short-lived. This “ratchet” effect limits the agents’ encouragement incentives, placing upper bounds on the equilibrium effort levels.⁴

We consider agents with heterogeneous costs of effort and show that agents with higher costs may work harder than the socially efficient level. Intuitively, the agents who are teamed with low-cost partners have a greater incentive to encourage them, since a low-cost agent’s effort is more sensitive to his beliefs. Therefore, the higher-cost agents, whose socially efficient effort levels are relatively low, may overwork. However, we derive an output sharing rule that eliminates effort overprovision and makes bounds on equilibrium effort coincide with their socially efficient levels. Furthermore, as the feedback becomes more responsive to manipulation, the equilibrium effort choices converge to socially efficient levels, provided that the right sharing rule is used.

Our theory has implications for important questions concerning team design, namely optimal timing of feedback as well as if and when it is desirable to use joint performance measures. To study this, we introduce a principal who controls the release of feedback to the agents and seeks to maximize total output. We show that the principal may find it optimal to restrict agents’ access to feedback. Doing this allows the principal to leverage the encouragement incentives by controlling the magnitude of the ratchet effect. For instance, suppose that the principal fully blocks access to feedback during some final phase. Then, during this phase, the ratchet effect will be eliminated as agents’ underperformance does not affect others’ beliefs. Therefore, any optimism created early on will be long-lived, strengthening the encouragement incentives. We provide conditions under which “one-time feedback” maximizes team output. We also consider the problem of a principal who can observe individual outputs of team members. We show that such a principal may nevertheless choose to reward agents on joint output. Doing so introduces free-riding incentives, but it also activates the encouragement effect. We show that under certain conditions, the latter can be sufficiently strengthened via restricted feedback to overcome the former, leading the principal to prefer joint performance measures.

Another contribution of this paper is to provide a framework that is tractable yet rich enough to study our economic question. Analyzing dynamic team incentives with uncertainty requires a model in which learning interacts with unobservable actions. In such models, however, characterization of behavior off the equilibrium path is severely complicated, as an agent’s deviation may cause her private belief to diverge from the public belief. Our setup circumvents this problem by separating the feedback from the production process, which greatly simplifies belief updating off the equilibrium path. Moreover, the speed of learning in our model does not depend on agents’ action, eliminating their incentives for experimentation. This distinguishes our model from the strategic experimentation literature (Bolton and Harris 1999, Keller et al. 2005,

⁴The ratchet effect—the effect of potentially causing high expectations of the agent’s future action—is extensively analyzed in the literature on dynamic agency models with asymmetric information (Weitzman 1980, Freixas et al. 1985) and dynamic moral hazard with learning and symmetric uncertainty (Bhaskar 2014, Prat and Jovanovic 2014, Cisternas 2018a, Bhaskar and Mailath 2019).

Bonatti and Hörner 2011). We believe that our framework opens further possibilities to analyze various aspects of team and feedback design.⁵

Our work connects to multiple strands of the literature. First, as we have noted, a large literature on teams analyzes moral hazard in groups (Olson 1965, Alchian and Demsetz 1972, Hölmstrom 1982). The literature mostly suggests that cooperation can be sustained by “punishments” of past behavior in the form of either lower monetary transfers or future non-cooperation by teammates.⁶ Our paper demonstrates that encouragement incentives in the presence of uncertainty could alleviate free-riding when contractual remedies are not available.

In a paper closely related to ours, Bonatti and Hörner (2011) consider dynamic moral hazard in teams in the presence of uncertainty. They utilize an exponential bandit framework in which the interaction ends when the common project has a “breakthrough.” The probability of a breakthrough depends on the agents’ instant effort level and the type of the project. By contrast, in our model, the total output of the project increases in agents’ cumulative effort. To the best of our knowledge, ours is the first paper to tractably incorporate learning in a team production model where agents’ payoffs depend on the whole history of effort.

Second, as noted earlier, incentives to manipulate others’ beliefs by attempting to influence realizations of noisy signals have been investigated in various contexts. Since Hölmstrom (1999), the literature on career concerns has analyzed the “signal-jamming” incentives of a manager who attempts to affect the market belief about his innate ability. Riordan (1985) (oligopoly) and Fudenberg and Tirole (1986) (entrant–incumbent game) examine a firm’s incentive to make the competing firm more pessimistic about future profitability. Recently, Cisternas (2018a) substantially expands the career concerns model by allowing general (nonlinear) payoffs for the long-run player in a stationary environment. He shows how the ratchet effect shapes the player’s equilibrium incentives.⁷ Our paper complements Cisternas (2018a) by analyzing the evolution of the ratchet effect in a nonstationary environment. In general, we contribute to the signal-jamming literature by analyzing such incentives in a team production environment.

Third, the economic literature has identified various forms of “encouragement effect” that manifest themselves via mechanisms that are different from ours. In games with complete information, the literature on dynamic contribution games (Admati and Perry 1991, Marx and Matthews 2000, Yildirim 2006) shows that a public project can be completed by agents who contribute small amounts from time to time.

⁵Our model can trivially accommodate heterogeneity among agents with respect to their productivity or ability to influence informative feedback and, less trivially, heterogeneity with respect to information. This, for instance, would open the door to addressing questions on allocation of heterogeneous agents into teams operating within an organization.

⁶In the literature on contracts with many agents, a group contract based on total output can mitigate moral hazard in teams (Hölmstrom 1982, Legros and Matthews 1993); in repeated partnership games, the threat of future non-cooperation following a deviation sustains various equilibrium dynamics (Radner et al. 1986).

⁷Cisternas (2018b) generalizes the career concerns model along another dimension by allowing investment in human capital.

Georgiadis (2015) provides a continuous-time framework and derives insights for various team design issues. These papers assume that the payoff is realized only when the project's state reaches a prespecified threshold. Therefore, the effort choices at different points in time are strategic complements, which is the channel through which the encouragement effect operates.⁸

Several papers analyze the encouragement effect in games with incomplete information, but differ from ours in one or more of these crucial aspects. We contribute to this literature by identifying a novel channel by which encouragement can operate. The encouragement mechanism in our paper is one of signal-jamming, and, thus, its existence crucially depends on (i) the presence of symmetric uncertainty, (ii) unobservable actions, and (iii) payoff externalities. Bolton and Harris (1999) consider a multi-agent experimentation model in which agents' effort choices are publicly observable and payoff is not shared. They show that the possibility of eliciting future experimentation by others encourages current experimentation. Dong (2018) analyzes how an encouragement effect can result from asymmetric information in a canonical exponential bandit model with observable effort choices. She shows that the better-informed player increases his effort to "signal" his optimism to the uninformed player, leading to an increase in both players' efforts.⁹ Campbell et al. (2014) consider a public good provision problem in which players decide whether to directly and credibly disclose their private information about production successes. They find that when the deadline is close, players hide successes to increase others' effort incentives.

Fourth, our result regarding the benefit of joint performance measures relates to several papers. Che and Yoo (2001) demonstrate that it may be desirable to use joint performance measures in dynamic environments when team members have an advantage in monitoring each other. Dai and Toikka (2018) reach an analogous conclusion in a static environment where there is large ambiguity about the production technology and the principal maximizes his payoff guarantee. In a contest model, Halac et al. (2017) show that changing both the reward and the information structure could improve the outcome. Our result that the use of joint performance measures can be optimal only when used in conjunction with information control is reminiscent of their conclusion. However their mechanism is distinct from ours: in their model, since the budget is fixed, agents who succeed early wish to *discourage* others' effort under the shared prize scheme and restricted information helps eliminate this adverse incentive.

Finally, there is a large literature in management regarding the effect of team potency—collective belief regarding the team's ability to be successful—on team performance (Mathieu et al. 2008). The literature finds that team potency has a positive impact on performance through their respective effects on the actions of the team members (Gully et al. 2002). Our paper contributes to the literature by suggesting the novel hypothesis that the team members' ability to affect team potency via feedback may have a positive effect on performance.

⁸See also Georgiadis (2017) for the effects of deadlines and monitoring frequencies on free-riding incentives.

⁹Klein and Wagner (2018) demonstrate that the encouragement effect can arise in a strategic investment model with experimentation due to private information.

The remainder of this paper is organized as follows. [Section 2](#) describes the model. [Section 3](#) characterizes the equilibrium and discusses its dynamics. [Section 4](#) conducts comparative statics. [Section 5](#) analyzes the effect of the feedback timing and performance measures on team incentives. [Section 6](#) concludes. The [Appendix](#) contains all the omitted proofs.

2. MODEL

A team of N agents undertakes a common project over a fixed time period $[0, T]$, where $T < \infty$. At the beginning of the game, nature draws a persistent state θ from a Gaussian distribution $\mathcal{N}(\mu_0, 1/\nu_0)$. At every time t , agents simultaneously choose private effort levels $a_i(t)$ ($i = 1, \dots, N$). Agent i 's effort has a quadratic flow cost of $c_i a_i(t)^2/2$. Agents do not discount the payoffs.

Agents observe neither the state nor the others' effort choices, but make inferences based on public feedback $Y(t)$. The feedback evolves according to a stochastic process

$$dY(t) = \left(\theta + \kappa \sum_{i=1}^N a_i(t) \right) dt + \frac{1}{\sqrt{\eta}} dW(t), \quad (1)$$

where $W(t)$ is a standard Brownian motion. Note that the drift of the above process is affected by both the unknown state (θ) and the agents' unobservable actions ($a_i(t)$). Here, κ measures the responsiveness of feedback on the agents' efforts and η controls the feedback precision.

We make a restriction on the feedback timing that the agents observe $Y(t)$ at intervals of time $\Delta > 0$. Formally, assume (without loss of generality) that T/Δ is an integer and that the agents observe $Y(t)$ only at $t = k\Delta$ ($k = 1, \dots, T/\Delta$). Note that with no discounting, the agents do not change their effort levels until the new feedback is observed. This effectively makes our model a discrete-time model. In [Sections 3 and 4](#), we mostly focus on the limit case of continuously observable feedback ($\Delta \rightarrow 0$). Our hybrid model—discrete feedback with continuous effort choices—is especially useful in analyzing the optimal feedback timing, which we address in [Section 5](#).¹⁰

At the end of the project—that is, at time T —an output P is realized. The amount of output depends on both the project state and the total effort:

$$P = \theta \int_0^T \sum_{i=1}^N a_i(t) dt.$$

The output is shared according to an exogenous sharing rule $s = (s_1, s_2, \dots, s_N)$, with $s_i \geq 0$ and $\sum_{i=1}^N s_i = 1$.

¹⁰Another advantage of our model, compared to a full continuous-time version, is that it admits a unique Nash equilibrium ([Theorem 1](#)). In [Appendix B](#), we show that the unique Nash equilibrium converges to the unique linear Markov perfect equilibrium of the continuous-time version of our model, providing justification of the Markov perfection criteria.

Given an effort profile $a = \{(a_1(t), \dots, a_N(t))\}_{t \in [0, T]}$ such that $a_i(t)$ is measurable with respect to the information available at time t , agent i 's expected payoff is

$$\mathbb{E}_{a, \theta} \left[s_i P - \int_0^T c_i \frac{a_i(t)^2}{2} dt \right] = \mathbb{E}_{a, \theta} \left[\int_0^T \left(s_i \theta \sum_{i=1}^N a_i(t) - c_i \frac{a_i(t)^2}{2} \right) dt \right].$$

A public history of length $t = k\Delta$ is a sequence $Y^k = \{Y(k'\Delta)\}_{\{k'=1, \dots, k\}}$. Agent i 's private history of length t , $h_i(t)$, is the combination of public history and his own past actions up to time t . Formally, $h_i(t) = \{Y^{\hat{k}(t)}, \{a_i(t')\}_{t' \in [0, t]}\}$, where $\hat{k}(t) = \max\{k : k\Delta \leq t\}$. A pure strategy of agent i is a mapping from his private histories into $\mathbb{R}$. We focus on pure strategy Nash equilibria. Note that because of the full support assumption on the feedback process, the only deviations that are detectable by agent i are his own. Consequently, all Nash equilibria are perfect Bayesian.

Remarks about the model Two aspects of our model deserve further elaboration. First, note that the output P exhibits complementarity between effort $a_i(t)$ and the state θ . Such complementarity ensures that the expected marginal product of an agent's effort is higher when he is more optimistic and thereby generates encouragement incentives via belief manipulation. Second, feedback $Y(t)$ is additively separable in agents' actions and the state. This specification is not necessary for the presence of encouragement incentives, but it renders our dynamic model tractable. In particular, as [Section 2.1](#) clarifies, this assumption implies that the speed of learning is independent of agents' actions and, thus, eliminates considerations of experimentation motives. This feature distinguishes our mechanism from those in the literature on strategic experimentation ([Bonatti and Hörner 2011](#), [Keller et al. 2005](#)).

2.1 Belief updating

In this subsection, we analyze the agents' belief updating process on and off the equilibrium path. As a benchmark case, suppose that the agents continuously observe feedback $Y(t)$ at every instant. Let $a^*(t) = \{a_i^*(\tau)\}_{i, \tau \in [0, t]}$ be the agents' common conjecture about their effort paths. Define

$$Z(t) = \frac{1}{t} \left(Y(t) - \kappa \int_0^t \sum_{i=1}^N a_i^*(t') dt' \right).$$

By (1), for any $t > 0$, $Z(t)$ is distributed normally with mean θ and precision ηt as long as the agents follow the conjectured effort path.

Define *public belief* as the common posterior belief under the assumption that all agents follow $a^*(t)$. Then public belief at time t is Gaussian with mean

$$\mu(t) \equiv \mu(Y(t), a^*(t)) = \frac{\nu_0 \mu_0 + \eta t Z(t)}{\nu_0 + \eta t} \quad (2)$$

and precision $\nu(t) = \nu_0 + \eta t$.¹¹

¹¹Strictly speaking, the public posterior mean and other posteriors are functions of relevant histories and not just time. To save on notation, we drop references to the specific history and simply index the beliefs by time whenever this leads to no confusion.

If an agent deviates from the conjectured effort path, however, his posterior belief differs from the public belief. Define agent i 's *private belief* at time t as a function of his private history $h_i(t)$ and his conjecture of others' effort paths $a_{-i}^*(t) = \{a_j^*(\tau)\}_{j \neq i, \tau \in [0, t]}$. Define

$$\hat{Z}_i(t) = \frac{1}{t} \left(Y(t) - \kappa \int_0^t \left(a_i(t') + \sum_{j \neq i} a_j^*(t') \right) dt' \right).$$

Then agent i 's private belief at time t is Gaussian with mean

$$\hat{\mu}_i(t) = \frac{\nu_0 \mu_0 + \eta t \hat{Z}_i(t)}{\nu_0 + \eta t} \quad (3)$$

and precision $\hat{\nu}_i(t) = \nu(t) = \nu_0 + \eta t$. Note that if all agents follow $a^*(t)$, private belief coincides with public belief (that is, $\hat{\mu}_i(t) = \mu(t)$) for any t .

For future use, we define $\rho(t) = \eta / (\nu_0 + \eta t)$, and express (2) and (3) as

$$\mu(t) = (1 - \rho(t)t) \mu_0 + \rho(t)t Z(t), \quad (4)$$

$$\hat{\mu}_i(t) = (1 - \rho(t)t) \mu_0 + \rho(t)t \hat{Z}_i(t). \quad (5)$$

Note that while agent i 's private belief is not affected by his own deviation from the conjectured path, the mean of public belief is. The parameter $\rho(t)$ captures the rate that an agent's effort at time t affects the future public belief.

In our model, feedback is not observed continuously, but is observed at intervals of time $\Delta > 0$. In this case, the belief updating processes are identical to (2) and (3), but beliefs are updated only at the moments when feedback is revealed. Formally, the public (private) belief path $\mu^\Delta(t)$ ($\hat{\mu}_i^\Delta(t)$) is given by $\hat{\mu}_i^\Delta(t) = \hat{\mu}_i(k\Delta)$ ($\mu^\Delta(t) = \mu(k\Delta)$) for $t \in [k\Delta, (k+1)\Delta)$.

3. EQUILIBRIUM

Our model admits a unique Nash equilibrium. This equilibrium has a remarkably simple structure: *After any history*, the equilibrium action of each agent is linear in the mean of his private posterior belief. Our main result, [Theorem 1](#), establishes the uniqueness of equilibrium and characterizes the equilibrium in the continuous-feedback limit (as $\Delta \rightarrow 0$) as solutions of a system of ordinary differential equations. We relegate all formal proofs to the [Appendix](#).

THEOREM 1. *For any $\Delta > 0$, there exists a unique Nash equilibrium. In equilibrium, agent i 's action is given by*

$$a_i^*(h_i(t); \Delta) = \xi_i^\Delta(t) \hat{\mu}_i^\Delta(t),$$

where the coefficient $\xi_i^\Delta(t)$ is a deterministic function of time.

Let $a_i^*(t)$ be agent i 's equilibrium action under the continuous-feedback limit ($\Delta \rightarrow 0$). Then $a_i^*(t) = \xi_i(t)\hat{\mu}_i(t)$, where $\{\xi_i\}_{i=1}^N$ satisfies

$$\dot{\xi}_i(t) = -\rho(t)\kappa \left[\frac{s_i}{c_i} \sum_{j \neq i} \xi_j(t) - \xi_i(t) \left(\xi_i(t) - \frac{s_i}{c_i} \right) \right] \quad (6)$$

for $i = 1, \dots, N$, along with terminal conditions $\xi_i(T) = s_i/c_i$.

We call the coefficient $\xi_i(t)$ agent i 's *belief sensitivity of effort* at time t : It measures the rate at which an agent responds to changes in his belief. The belief sensitivity of effort is deterministic and varies only with calendar time. Next, we provide a heuristic derivation of (6) that characterizes the equilibrium in the continuous-feedback limit.^{12,13}

First, note that because of our linear-quadratic structure, the equilibrium effort level is linear in the marginal return to effort. In our model, there are two sources of return to effort. First is the direct return due to its contribution to the output. If θ is commonly known, this is the only source of return.

The presence of uncertainty creates a second source of return, as agent i 's behavior affects others' future effort levels. To understand this, suppose that agent i chooses $a_i(t) = a_i^*(t) + \alpha$ over a small interval $[t, t + dt)$ for some $\alpha > 0$, while the others follow the equilibrium profile. Then (1) implies that for any $t' \geq t + dt$, agent i 's deviation boosts feedback $Y(t')$ by $\kappa\alpha dt$. Since her deviation is not detected, other agents form a more optimistic belief than agent i does: (5) implies that the positive change in others' belief is $\rho(t')\kappa\alpha dt$. Then following the equilibrium strategies, other agents increase their effort levels at $t' \geq t + dt$ by

$$\rho(t')\kappa\alpha dt \sum_{j \neq i} \xi_j(t'). \quad (7)$$

To calculate the correct value of the belief manipulation return, however, we also need to consider the effect of belief divergence after $t + dt$. Recall that after an upward deviation at time t , the other agents' belief is higher than that of agent i by $\rho(t)\kappa\alpha dt$. Given their optimistic belief, other agents form a high expectation about agent i 's performance, but agent i optimally chooses a lower effort given her pessimistic belief. Specifically, agent i 's effort level for $[t + dt, t + 2dt)$ falls short of others' expectations by $\xi_i(t + dt)\rho(t)\kappa\alpha dt$. Combining with (7), the net effect of initial deviation at t and the

¹²In the [Appendix](#), we derive a closed-form formula for the unique Nash equilibrium for arbitrary $\Delta > 0$. With $\Delta > 0$, an agent's incentives remain unchanged and, thus, his equilibrium effort must be constant over the interval $[k\Delta, (k+1)\Delta)$. This observation enables us to treat the game as a discrete-time game with T/Δ "periods." Our equilibrium construction and the uniqueness argument utilize backward induction, where the linear-quadratic structure implies that each agent has a unique best response after any history, which is linear in his private posterior mean.

¹³See [Appendix B](#) for a characterization of the unique linear Markov equilibrium for a continuous-time model in which the feedback $Y(t)$ is continuously observed. Also, [Appendix B](#) shows that this equilibrium coincides with the continuous-feedback limit of the equilibrium in [Theorem 1](#) characterized by (6).

resulting belief wedge at $t + dt$ on the others' effort level at $t' \geq t + 2dt$ is

$$\underbrace{\left[1 - \xi_i(t + dt)\rho(t)\right]}_{\text{decay due to the ratchet effect}} \rho(t')\kappa\alpha dt \sum_{j \neq i} \xi_j(t'). \quad (8)$$

Comparing (7) with (8) implies that the effect of the belief divergence, which we refer to as the *ratchet effect*, dampens the marginal return from the belief manipulation effect.

Given our understanding of the belief manipulation effect, we now turn to the derivation of (6) by comparing the values of $\xi(t)$ and $\xi(t + dt)$. Since the direct return of effort is constant over time, the only difference between $\xi(t)$ and $\xi(t + dt)$ is about return from belief manipulation. This difference is expressed by the recursive formulation

$$\begin{aligned} \xi_i(t) - \xi_i(t + dt) = & \underbrace{\frac{s_i}{c_i}\rho(t + dt)\kappa dt \sum_{j \neq i} \xi_j(t + dt)}_{\text{returns from manipulating } (t + dt) \text{ belief}} \\ & - \underbrace{\rho(t + dt)\kappa dt \xi_i(t + dt)}_{\text{extra discounting}} \underbrace{\left(\xi_i(t + dt) - \frac{s_i}{c_i}\right)}_{\text{future incentives at } t + dt} . \end{aligned}$$

ratchet effect

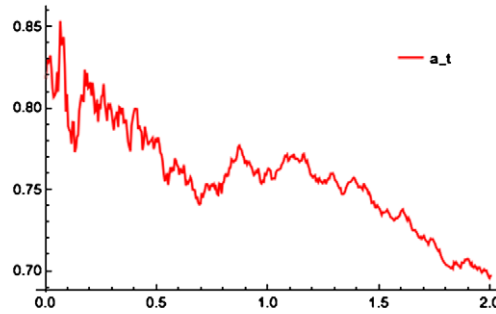
This formulation highlights two differences between belief manipulation incentives over $[t, t + dt)$ versus over $[t + dt, t + 2dt)$. First, over $[t, t + dt)$, agent i can influence others' efforts over $[t + dt, t + 2dt)$, while there is no such opportunity over $[t + dt, t + 2dt)$: This difference corresponds to the first term on the right-hand side. Second, an increase in effort at time t creates a wedge between public and private belief at $t + dt$, which leads to an extra period of ratchet effect. This is manifested as additional discounting on the future returns from time $t + dt$ onward.¹⁴ Simplifying and rearranging the above expression yields (6).

The system (6) admits a unique solution characterized by

$$\xi_i(t) = \frac{s_i}{c_i} \left(1 + \rho(t)\kappa \int_t^T \sum_{j \neq i} \xi_j(x) e^{-\int_t^x \rho(l) dl} e^{-\int_t^x \kappa \rho(l) \xi_i(l) dl} dx \right). \quad (9)$$

This characterization has an intuitive interpretation. Since the cost of effort is quadratic, after any history, agent i 's optimal effort is equal to her marginal private return per unit of effort cost. The first term in (7) captures the direct marginal return on effort, which is given by $(s_i/c_i)\hat{\mu}_i(t)$. The second term captures the indirect benefit due to encouragement via belief manipulation. The indirect benefit can be split into two parts: the initial boost in public belief given by $\rho(t)\kappa$ and the rate at which this boost decays over time. The decay of the belief boost comes from the two sources. First, as more information is

¹⁴The form of this additional discounting can be better understood by referring to (8). At time $t + \Delta$, agent i 's effort choice falls short of others' expectations and, therefore, the belief boost created at time t is further discounted. Note that $\xi_i(t + dt)$ represents the full marginal return on effort while s_i/c_i represents the direct productive benefit. Therefore, $\xi_i(t + dt) - s_i/c_i$ corresponds to the future returns.

FIGURE 1. Simulated equilibrium effort path ($N = 5$, $\kappa, \nu = 1$, $T = 2$).

revealed in the future, the public belief converges to the unbiased belief; this is captured by the first exponential term ($e^{-\int_t^x \rho(l) dl}$). The second exponential term ($e^{-\int_t^x \kappa \rho(l) \xi_i(l) dl}$) captures the ratchet effect: For every instant after the deviation, agent i underperforms relative to others' expectations, speeding up the convergence of the public belief to the unbiased belief.

3.1 Equilibrium dynamics

Figure 1 depicts a simulated path of the equilibrium effort level $a^*(t) = \xi(t)\mu(t)$ in the case of symmetric agents. Note that the agents' effort paths are stochastic and typically non-monotonic, as agents' beliefs follow the realizations of the Gaussian noise in the feedback. On the contrary, the belief sensitivity of effort $\xi_i(t)$ is deterministic, providing a clearer insight on the equilibrium dynamics. Moreover, ex ante expected output of a team is a linear function of the sum of agents' belief sensitivities.¹⁵ Therefore, in the rest of the paper, we mostly focus on analyzing the characteristics of $\xi_i(t)$.

The next proposition establishes the monotonicity of $\xi_i(t)$ over time and characterizes its lower and upper bounds.

PROPOSITION 1. For any $i = 1, \dots, N$,

- (i) $\xi_i(t)$ is strictly decreasing for all $t \in [0, T)$, with $\xi_i(T) = s_i/c_i$,
- (ii) $\xi_i(t) < \bar{\xi}_i$ for any $t \in [0, T)$, where $\bar{\xi}_i = \sqrt{\frac{s_i}{c_i}} \sum_{j=1}^N \sqrt{\frac{s_j}{c_j}}$.

Recall from (6) that $\xi_i(t)$ consists of two sources of marginal return to the agent's effort: the direct return and the indirect return from belief manipulation. Since the direct return is constant over time (s_i/c_i for all t), the dynamics of $\xi_i(t)$ shows the evolution of the belief manipulation incentives over time.

¹⁵Since $a_i^*(t) = \xi_i(t)\hat{\mu}_i(t)$ and $\hat{\mu}_i(t) = \mu(t)$ on the equilibrium path, the ex ante expected payoff of a team is given by

$$\mathbb{E}_{t=0} \left[\int_0^T \theta \sum_{i=1}^N a_i^*(t) dt \right] = \int_0^T \mathbb{E}_{t=0} [\theta \mu(t)] \sum_{i=1}^N \xi_i(t) dt = \int_0^T \left(\mu_0^2 + \frac{\eta t}{\nu_0(\nu_0 + \eta t)} \right) \sum_{i=1}^N \xi_i(t) dt.$$

The first part of [Proposition 1](#) implies that the return from belief manipulation monotonically decreases over time and then disappears at the end of the project. There are two reasons for this decline. First, higher t simply means that the project is closer to the deadline, reducing the returns from encouraging others' future effort. Second, as the agents learn θ more precisely over time, they place a smaller weight on new feedback in updating their beliefs (i.e., $\rho(t)$ declines), making it costlier to manipulate beliefs.

The second part implies that there is a limit on the size of the return from belief manipulation. This is a direct manifestation of how the ratchet effect dampens the belief manipulation incentives. Notice that in (6), (the change in) the gain from the belief manipulation (the first additive term) is linear in $\xi_j(t)$, but (the change in) the ratchet effect (the second additive term) is quadratic in $\xi_i(t)$. Therefore, as belief manipulation incentives increase, the ratchet effect catches up to the linear gains. The upper bound of $\xi_i(t)$ is reached when the ratchet effect fully washes out the direct belief manipulation incentives.¹⁶

[Proposition 1](#) holds only in the continuous-feedback limit ($\Delta \rightarrow 0$). If feedback is observed in discrete intervals, the evolution of the ratchet effect may induce non-monotonic $\xi_i(t)$. For example, suppose that $\xi_i(t)$ is very large (larger than $\bar{\xi}_i$ in [Proposition 1](#)). Then a small increase in an agent's effort at $t - \Delta$ leads to a large discrepancy between his effort level at t and others' expectation thereof. This in turn implies a large decrease in others' belief at $t + \Delta$, rendering the ratchet effect at $t - \Delta$ very severe. This can lead to weaker effort incentives at $t - \Delta$ than at t . For small Δ , the ratchet effect kicks in immediately so that $\xi_i(t)$ can never exceed $\bar{\xi}_i$, eliminating this possibility.¹⁷ This observation also hints at our subsequent analysis in [Section 5](#), which demonstrates how a principal, via use of information control, can induce $\xi_i(t)$ to exceed the bound $\bar{\xi}_i$.

3.2 Role of project uncertainty

Risk-taking is considered to be one of the main elements of entrepreneurial behavior. The literature suggests various explanations for risk-taking behavior, such as the higher premiums or risk-loving preferences of entrepreneurs. This paper identifies an alternative motivation: *Undertaking an uncertain project can benefit organizations by mitigating the free-rider problem.*

It is easy to show that if θ is known (i.e., $\nu_0 = \infty$), the belief manipulation effect disappears and the agents exert their myopic optimum effort ($s_i\theta/c_i$) at any t . The following proposition formalizes this and analyzes the impact of ν_0 on effort incentives away from this limit.

PROPOSITION 2. *For any $t \in [0, T)$ and $i = 1, \dots, N$, the belief sensitivity of effort $\xi_i(t)$ decreases in ν_0 and converges pointwise to s_i/c_i as $\nu_0 \rightarrow \infty$.*

¹⁶The existence of an upper bound is in contrast to the outcome in the model in [Hölmstrom \(1999\)](#), where the optimal effort level of the long-run player may diverge as the horizon becomes longer. In [Hölmstrom \(1999\)](#), the agent's marginal product of effort is independent of beliefs; therefore, even when the public belief diverges from his own private belief, no wedge appears between his effort choice and the public expectation, whereas such a wedge is the driver of the ratchet effect in our model.

¹⁷We further elaborate on this possibility in [Appendix C](#).

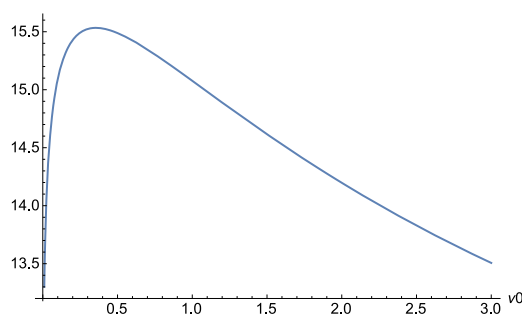


FIGURE 2. Ex ante payoff as a function of ν_0 ($T = 10$, $N = 5$, $\kappa = \eta = 1$, $c_i = 1$, $s_i = 1/5$).

Another necessary condition for the existence of belief manipulation incentives is imperfect monitoring. If effort choices are perfectly observable, agents cannot bias each other's beliefs upward by boosting the feedback and, thus, in any equilibrium, they choose their myopic optimal effort $s_i \mu(t)/c_i$. In the team production literature, the inability or a high cost to monitor individual effort has mostly been considered as an obstacle to inducing cooperation, as it limits the ability to accurately punish undesirable actions.¹⁸ Our result implies that in the presence of uncertainty, imperfect monitoring may generate another strategic effect that helps to overcome the free-riding problem.

An important implication of Proposition 2 is that the presence of uncertainty may increase the welfare of the agents. However, the standard cost of uncertainty—arising from agents' uninformed decisions—still exists in this model. This argument leads to a question regarding the optimal level of project uncertainty. To address this, we consider classes of projects indexed by $(\mu_0; \nu_0)$ that would deliver the same expected total payoff to the team under complete information, i.e., if the uncertainty were to resolve prior to the production stage. Figure 2 plots the expected welfare of the agents while varying the prior precision ν_0 within this class. In Appendix D, we conduct a formal and complete analysis of Figure 2.¹⁹ The graph shows that there exists an interior ν_0 that optimally balances the trade-off between benefit from the belief manipulation effect and the cost from uninformed decisions.

4. COMPARATIVE STATICS

The tractability of our environment allows for several comparative statics results that add to the understanding of the belief manipulation incentives. The next propositions

¹⁸As Alchian and Demsetz (1972) write, "... In team production, marginal products of cooperative team members are not so directly and separably (i.e., cheaply) observable. What a team offers to the market can be taken as the marginal product of the team but not of the team members. The costs of metering or ascertaining the marginal products of the team's members is what calls forth new organizations and procedures."

¹⁹The team's surplus under complete information is a convex function of true project quality (θ), which creates an inherent value for initial uncertainty even without its role in boosting effort incentives. To isolate the impact of the latter, when calculating the impact of decreased ν_0 on total payoff, we reduce μ_0 just enough to nullify the first effect.

collect these results. Again, our results focus on the belief sensitivity of effort $\xi_i(t)$ in the continuous-feedback limit ($\Delta \rightarrow 0$), unless otherwise noted.

PROPOSITION 3. *For any $t \in [0, T)$ and $i = 1, \dots, N$, belief sensitivity of effort $\xi_i(t)$*

- (i) increases in κ and converges pointwise to $\bar{\xi}_i$ as $\kappa \rightarrow \infty$,*
- (ii) increases in η ,*
- (iii) decreases in c_i as well as c_j ($j \neq i$).*

In addition, when the agents are homogeneous and the sharing rule is symmetric, for any $t \in [0, T)$, $\xi(t)$ decreases, but $N\xi(t)$ increases in N .

Intuitively, returns to belief manipulation are greater as feedback becomes more sensitive to effort (higher κ) or feedback becomes more precise (higher η), because in either case, feedback carries a larger weight in the updating process (4). The convergence result in Proposition 3(i) shows that as κ becomes arbitrarily large, the counteracting ratchet effect cancels out the direct belief manipulation effect, leading to the upper limit on the belief sensitivity specified in Proposition 1.

It is straightforward that agent i 's incentive decreases in her own cost parameter (c_i). Perhaps more interestingly, her belief sensitivity also decreases in the others' cost parameters (c_j). This highlights the *endogenous strategic complementarity* generated by the belief manipulation incentives. Under a smaller c_j , agent j 's incentives for effort directly increase. Then agent j 's effort choice will be more sensitive to her belief, which in turn increases agent i 's returns from his own effort.

The last part of Proposition 3 refers to the impact of team size on the effort incentives. Intuitively, increasing the size of the team strengthens both the free-riding incentives and the belief manipulation incentives. The last part of Proposition 3 implies that the former negative effect dominates the latter as N grows large. However, this weakening of individual effort incentives is more than compensated by the increase in the team size, leading to higher total effort incentives.²⁰ This is in contrast to the complete information counterpart, where the sum of belief sensitivities $N\xi(t) = N \cdot (1/cN) = 1/c$ stays constant.

4.1 Sharing rule and convergence to the first-best

That agents' effort incentives are stronger when others' effort is more sensitive to belief is highlighted by Proposition 3(iii). In fact, these incentives can be too strong relative to first-best if an agent has high effort cost, but his teammates have low costs and, thus, are very sensitive to beliefs. Figure 3 plots the belief sensitivities in a two-agent team with asymmetric effort costs ($c_1 = 1$, $c_2 = 0.6$), with the dashed lines corresponding to

²⁰Bolton and Harris (1999), in a multi-agent experimentation model, show that while per capita experimentation may decrease as the number of agents increases, this decrease is offset by the increase in the number of players, leading to a nondecreasing value function.

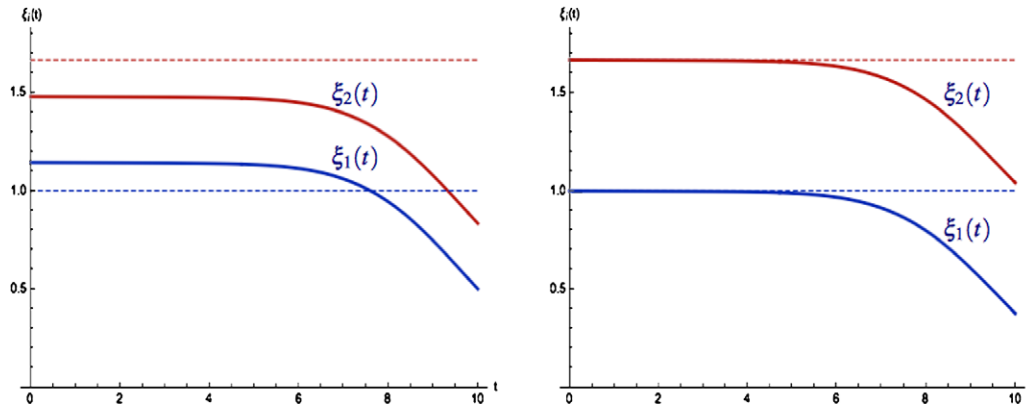


FIGURE 3. Asymmetric effort costs $c_1 = 1$ and $c_2 = 0.6$: left panel, $s = (1/2, 1/2)$; right panel, $s = (3/8, 5/8)$.

the socially efficient levels $1/c_i$.²¹ In the left panel, where the agents' shares are equal, the high-cost agent's effort exceeds its socially efficient level ($1/c_1$), while the low-cost agent's effort level is bounded away from $1/c_2$. Proposition 4 states that it is always possible to fine-tune the sharing rule such that the upper bound of $\xi_i(t)$ coincides with the socially efficient levels. The right panel of Figure 3 demonstrates this result by describing the equilibrium $\xi_i(t)$ under $s_i^* = (3/8, 5/8)$.

PROPOSITION 4. *If the sharing rule $s = (s_1, s_2, \dots, s_N)$ satisfies*

$$s_i = s_i^* \equiv \frac{\frac{1}{c_i}}{\sum_{j=1}^N \frac{1}{c_j}},$$

the upper bound of $\xi_i(t)$ coincides with its socially efficient level, that is, $\bar{\xi}_i = 1/c_i$.

Note that s_i^* in Proposition 4 is not necessarily the share structure that maximizes the expected payoff, as $\xi_i(t)$ is generally bounded away from $\bar{\xi}_i$. In fact, $\xi_i(t)$ may not reach its upper bound for any t , as shown in Proposition 2. However, the following corollary, which is a straightforward implication of Propositions 3 and 4, identifies a limit case under which the equilibrium $\xi_i(t)$ converges to the socially efficient level.

COROLLARY 1 (Convergence to socially efficient outcome). *If the sharing rule satisfies $s_i = s_i^*$ for all i , the belief sensitivity $\xi_i(t)$ converges pointwise to its socially efficient level for $t \in [0, T)$ as $\kappa \rightarrow \infty$.*

²¹ Given the uncertainty about θ and the absence of learning considerations, the socially optimal level of effort for each agent i would be $\mu(t)/c_i$, where $\mu(t)$ is the unbiased mean belief about θ given the realizations of the feedback. Since learning happens exogenously in our model, there are no intertemporal trade-offs between current and future surplus created by varying learning speeds. Therefore, the socially optimal efforts simply maximize the current surplus given the current information. Therefore, the socially optimal outcome would prescribe a constant belief sensitivity of $\xi_i^* = 1/c_i$.

To understand this seemingly unexpected result, consider the marginal return to agent i 's effort when all agents' belief sensitivities are close to their upper bounds (that is, $\xi_j(t) = \bar{\xi}_j$ for all j) and κ is arbitrarily large. Equation (9) implies that if agent i marginally increases her effort, then the belief manipulation effect provides an initial boost to the others' effort levels, the size of which is approximately $\rho(t)\kappa \sum_{j \neq i} \bar{\xi}_j$. Then, due to the ratchet effect, the initial boost decays at an approximate rate of $\rho(t)\kappa \bar{\xi}_i$.²² Therefore, the total present value of the increase in the others' efforts is approximately $(\sum_{j \neq i} \bar{\xi}_j)/\bar{\xi}_i$. This is simply the ratio of the sum of others' belief sensitivities (which determine the size of the initial boost) and own belief sensitivity (which acts as a discount rate).

Adding the direct marginal product of agent i 's effort (1) and multiplying by agent i 's share (s_i) yields the private marginal return of agent i 's effort:

$$s_i \left(1 + \frac{\sum_{j \neq i} \bar{\xi}_j}{\bar{\xi}_i} \right) = s_i \frac{\sum_{j=1}^N \bar{\xi}_j}{\bar{\xi}_i}.$$

Social efficiency requires the private marginal return to equal the social marginal return, which implies that the sharing rule must satisfy

$$s_i \frac{\sum_{j=1}^N \bar{\xi}_j}{\bar{\xi}_i} = 1 \implies s_i = \frac{\bar{\xi}_i}{\sum_{j=1}^N \bar{\xi}_j}.$$

This sharing rule is clearly feasible and, indeed, the shares characterized in [Proposition 4](#) are found by plugging $\bar{\xi}_i = \frac{1}{c_i}$ into the above expression for s_i .

5. RESTRICTED FEEDBACK AND JOINT PERFORMANCE MEASURES

This section has two messages. First, a principal who can control the timing of the feedback can increase the team's output relative to its level in our unique equilibrium. Second, this improvement can be so drastic that an output-maximizing principal may be better off choosing to employ such a measure and reward the agents on joint output even when he is able to observe the individual outputs of the team members, in spite of the fact that the latter would eliminate all free-riding concerns.

²²By the time $x > t$, this boost is discounted for two reasons: (i) exogenous learning, by $e^{-\int_t^x \rho(l) dl}$, and (ii) ratcheting, by $e^{-\int_t^x \kappa \rho(l) \bar{\xi}_i dl}$. When κ is large, the variation in $\rho(l)$ is negligible compared to κ and, therefore, a linear approximation of the integral $\int_t^x \rho(l) dl$ by $\rho(t)(x - t)$ becomes appropriate, leading to this assertion.

5.1 Restricted feedback

Consider a principal whose payoff is an increasing function of the team's total output and who privately observes $Y(t)$. Before production begins he can commit to a “feedback schedule,” which is a subset $\mathbb{T} \subset [0, T]$, with the understanding that at each $t \in \mathbb{T}$, the principal publicly and credibly reveals $Y(t)$. In all the variations we consider below, agents' equilibrium efforts are linear in their mean beliefs so that their equilibrium strategies, as usual, are characterized by their deterministic belief sensitivities. Let $\xi_i^{\mathbb{T}}(t)$ represent the belief sensitivity of effort under feedback schedule $\mathbb{T}$. Then since agent i 's time- t effort level on the equilibrium path is $\xi_i^{\mathbb{T}}(t)\mu(t)$, the expected output of a team is given by

$$\begin{aligned} P(\mathbb{T}) &= \mathbb{E}_{t=0} \left[\int_0^T \theta \mu(t) \sum_{i=1}^N \xi_i^{\mathbb{T}}(t) dt \right] \\ &= \int_0^T \left(\mu_0^2 + \frac{\gamma_{\mathbb{T}}(t)}{\nu_0} \right) \sum_{i=1}^N \xi_i^{\mathbb{T}}(t) dt, \end{aligned} \quad (10)$$

where $\gamma_{\mathbb{T}}(t)$ is a measure of the precision of information that the agents have at time t . Specifically,

$$\gamma_{\mathbb{T}}(t) = \frac{\eta \tau_{\mathbb{T}}(t)}{\nu_0 + \eta \tau_{\mathbb{T}}(t)},$$

where $\tau_{\mathbb{T}}(t) = \sup\{t' \in \mathbb{T} | t' \leq t\}$ is the most recent date of feedback preceding time t .

The form of expected output reveals that, all else being equal, the principal prefers to give as much information as possible (i.e., the output is increasing in $\gamma_{\mathbb{T}}(t)$). Then the only potential reason he would delay release of feedback is its strategic impact on $\xi_i^{\mathbb{T}}(\cdot)$. Note that the expected output is also increasing in $\xi_i^{\mathbb{T}}$. Next we consider an example illustrating the channel by which such a restriction may help.

A numerical example To fix ideas, consider the simplest possible case: a team consisting of two symmetric agents, with the equal sharing rule. Furthermore, assume that $c_i = \eta = \nu_0 = \mu_0 = 1$. First, under “continuous feedback” ($\mathbb{T} = [0, T]$), the equilibrium belief sensitivities are calculated from (6) as

$$\xi_1(t) = \xi_2(t) = \frac{1}{1 + \left(\frac{1+t}{1+T} \right)^{\kappa}}.$$

We immediately verify that belief sensitivities are increasing in κ but are bounded. To illustrate how holding back feedback may boost belief sensitivities, consider a very crude scheme that reveals feedback at a unique instant, say $\tilde{t}$; i.e., consider $\mathbb{T} = \{\tilde{t}\}$. The belief sensitivities under this feedback schedule are given by

$$\xi_i(t) = \begin{cases} \frac{1}{2} \left[1 + \frac{1}{2} (T - \tilde{t}) \rho(\tilde{t}) \kappa \right] & \text{if } t < \tilde{t}, \\ \frac{1}{2} & \text{if } t \geq \tilde{t}. \end{cases}$$

Since there is no further feedback after $\tilde{t}$, the belief sensitivities are equal to the direct marginal benefit of effort (in this case, $1/2$). The expression for $\xi_i(t)$ for $t < \tilde{t}$ also accounts for this while additionally capturing the familiar benefit from belief manipulation. The expression for the latter can be understood as follows. The rate at which an agent's effort at $t \in [0, \tilde{t})$ boosts his teammate's time $\tilde{t}$ belief is $\kappa\rho(\tilde{t})$. Since the belief sensitivity of effort for $t > \tilde{t}$ is $1/2$, the initial boost in the teammate's effort is $\kappa\rho(\tilde{t})/2$. Importantly, lack of further feedback eliminates both further learning and any ratchet behavior, allowing this boost to remain constant over $t \in [\tilde{t}, T]$. This leads to a total effort boost of $(T - \tilde{t})\kappa\rho(\tilde{t})/2$, accounting for the second additive term in the expression.

The only gain from restricting the feedback in this manner is the boost in belief sensitivities prior to $\tilde{t}$.²³ This gain operates through elimination of the ratchet effect over $[0, \tilde{t})$. Recall that in our baseline model, the ratchet effect dampens earlier effort incentives because it leads to rapid decay of the optimism generated by boosts to feedback. By shutting down learning after $\tilde{t}$, the one-time feedback schedule completely eliminates the ratchet effect. Importantly, in our baseline model, the ratchet effect is most severe when κ is large and is the force that bounds belief sensitivities. Not surprisingly, and as is clear by inspecting the expression for $\xi_i(t)$, when the ratchet effect is shut down, as κ grows large, the earlier belief sensitivities grow without bound. This advantage of one-time feedback eventually becomes sufficient so that it is preferred by the principal when κ is large.

Optimality of one-time feedback schedules The following proposition shows that the forces at work in the above example exist under a general environment, so that one-time feedback induces a greater output than continuous feedback if κ is large. The proposition also characterizes the optimal timing of one-time feedback schedules.

PROPOSITION 5. *Among all one-time feedback schedules, expected output is maximized when feedback is given at*

$$t^* = \frac{\sqrt{\eta\nu_0 T + \nu_0^2} - \nu_0}{\eta}.$$

Moreover, there exists $\bar{\kappa}$ such that if $\kappa > \bar{\kappa}$, then the expected output under $\mathbb{T}^ \equiv \{t^*\}$ exceeds expected output under continuous feedback $\mathbb{T}^{**} \equiv [0, T]$.*

As described in the above example, the negative ratchet effect does not exist under one-time feedback, since there is no subsequent feedback after belief divergence. Proposition 5 implies that when κ is large enough, cost from the ratchet effect also becomes large so that the one-time feedback dominates the continuous feedback schedule. The optimal timing t^* of feedback balances the following trade-off. While delaying the feedback means a longer period for agents to work harder because of the belief manipulation incentives, the incentives become weaker as there exists a shorter period during which manipulated beliefs influence behavior.

²³When feedback is restricted in this manner, $\gamma_{\mathbb{T}}(t)$ is weakly smaller for all t and, for $t > \tilde{t}$, the belief sensitivities are also smaller than under continuous (unrestricted) feedback.

Next we show that the one-time feedback schedule is optimal for the principal among the class of “coarse” feedback schemes under sufficiently large κ . Let $\underline{\Delta}(\mathbb{T})$ represent the “coarseness” of a feedback schedule $\mathbb{T}$ defined by

$$\underline{\Delta}(\mathbb{T}) = \inf\{|t - t'| \mid t, t' \in \mathbb{T}\}.$$

Also, define $\mathcal{T}(\Delta)$ as a set of feedback schedules $\mathbb{T}$ with $\underline{\Delta}(\mathbb{T}) > \Delta$.

PROPOSITION 6. *Fix $\Delta > 0$. Then there exists $\bar{\kappa}$ such that for any $\kappa > \bar{\kappa}$, the one-time feedback schedule $\mathbb{T}^*$ defined in Proposition 5 induces the largest expected output among the feedback schedules in $\mathcal{T}(\Delta)$.*

When away from the continuous-feedback limit, i.e., if period length $\Delta > 0$, the coarseness of available feedback schemes is bounded below by Δ . Then Proposition 6 implies that for fixed period length Δ , as κ becomes sufficiently large, $\mathbb{T}^*$ is optimal among all feasible feedback schemes.²⁴

5.2 Joint versus individual performance measures

Now suppose that the principal observes the individual outputs of team members. By paying each agent his/her own production, the principal can completely eliminate free-riding. However, doing so also eliminates the belief manipulation incentives as agents do not gain from others' hard work.

Under this “individual performance measure” (IPM), the belief sensitivity of each agent is simply $1/c_i$ and, in particular, is independent of the feedback schedule. Therefore, the following result immediately follows by inspecting (10).

PROPOSITION 7. *Under the IPM, the output-maximizing feedback schedule is the continuous feedback $\mathbb{T}^{**}$.*

In light of Proposition 7, to argue that the principal may prefer joint performance measures (JPM, i.e., splitting the total output among agents) over IPM, it suffices to show that JPM together with the best one-time feedback schedule $\mathbb{T}^*$ is superior to IPM together with the continuous-feedback schedule $\mathbb{T}^{**}$. Figure 4 illustrates the equilibrium belief sensitivities under each of these combinations as well as the benchmark case of JPM with $\mathbb{T}^{**}$ for a symmetric team of two agents.

Note that under full-information feedback, IPM generates larger belief sensitivities and, therefore, larger expected output than JPM.²⁵ However, since under a one-time feedback schedule, the belief sensitivities over $[0, t^*)$ can be arbitrarily large, the inequality can be reversed. The following proposition provides sufficient conditions for this reversal.

²⁴In the limit when $\Delta \rightarrow 0$, “coarseness” rules out continuous feedback over any positive-length interval.

²⁵In general, this follows because under joint performance measures, the belief sensitivities are bounded above by $\xi_i < 1/c_i$.

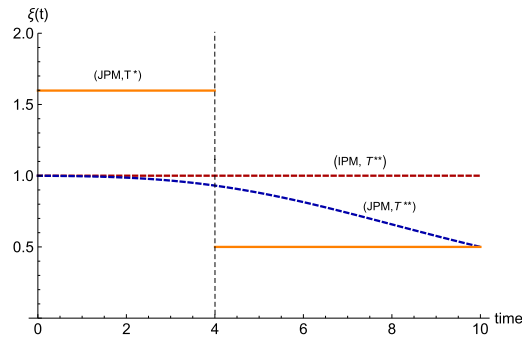


FIGURE 4. Equilibrium belief sensitivity under different feedback schemes and performance measures. ($T = 10$, $N = 2$, $c_i = 1$, $s_i = 1/2$.)

PROPOSITION 8. *There exists $\bar{\kappa}$ such that if $\kappa > \bar{\kappa}$, then the joint performance measure combined with one-time feedback generates more expected output than the individual performance measure with continuous feedback.*

The comparison of the two schemes $(\text{JPM}, \mathbb{T}^*)$ versus $(\text{IPM}, \mathbb{T}^{**})$ boils down to the relative impact on output of the additional early effort generated by the former versus the increased precision generated by the latter. This result is intuitive: By earlier arguments, the larger is the κ , the more advantageous is restricted feedback along with JPM, and this advantage grows unboundedly, whereas the advantage generated by eliminating free-riding is bounded.²⁶

6. CONCLUDING REMARKS

This paper contributes to the recent literature that investigates the effect of uncertainty on team behavior. While our model is simple, it highlights several interesting trade-offs that occur when joint production takes place over time and under uncertainty. Nevertheless, it may be of interest to understand how these trade-offs would interact with others that may appear in more general environments.

Our model can trivially accommodate further heterogeneity (e.g., with respect to agent productivity, ability to impact feedback) and discounting as long as the linear-quadratic-Gaussian structure is preserved. Moving away from the linear-quadratic-Gaussian framework appears to be a worthy path for future research even though it is more challenging. Allowing for more general production functions to accommodate interaction between agents' effort choices or a different relationship between the state and marginal product of effort may lead to strengthening or weakening of the encouragement effect that we are highlighting. Investigating these relationships can lead to

²⁶Che and Yoo (2001) and Dai and Toikka (2018) reach conclusions that are similar in spirit. Che and Yoo (2001) demonstrate that it may be desirable to use joint performance measures in dynamic environments when team members have an advantage in monitoring each other, while Dai and Toikka (2018) reach an analogous conclusion in a static environment where there is large ambiguity about the production technology and the principal maximizes his payoff guarantee.

fruitful conclusions about the optimality of various feedback schemes and performance measures in different technological environments.

A different type of heterogeneity that may be of interest is with respect to the agents' information. Specifically, our linear-quadratic-Gaussian model can be extended to accommodate the asymmetric information among the agents: some members (experts) may have more precise information about the state than others (novices), leading essentially to a dynamic version of [Hermalin \(1998\)](#). This extension may provide insights into optimal dynamic sharing rules between experts and novices.

The encouragement effect we highlight in this paper is likely to be present under different specifications. Here we assume that agents observe public feedback while their actions are private, and that the encouragement effect manifests itself as a form of signal jamming. Under the alternative specification where agents receive private feedback but take observable actions, effort incentives are likely to be boosted in an analogous fashion due to the agents' incentives to signal their information to others.²⁷ More broadly, analysis of the optimal contract combined with information design remains an interesting path for future research.

APPENDIX

The appendix consists of four parts. [Appendix A](#) contains all omitted proofs. [Appendix B](#) characterizes the linear Markov equilibrium in the continuous-time version of our model and shows the convergence result of the equilibrium in [Theorem 1](#) as $\Delta \rightarrow 0$. [Appendix C](#) discusses potential non-monotonicity of the equilibrium belief sensitivity when Δ is bounded away from 0. [Appendix D](#) describes an exercise that illustrates the trade-offs associated with determining the optimal level of uncertainty.

APPENDIX A: OMITTED PROOFS

Proof of Theorem 1

We prove [Theorem 1](#) in two steps. First, we show that there exists a unique Nash equilibrium in our model for any $\Delta > 0$ of the feedback interval ([Proposition 9](#)). Not only do we show the existence and uniqueness, but we also characterize the belief sensitivities ξ_i in recursive form ([12](#)). Second, we take the limit of $\Delta \rightarrow 0$ and show that the recursive forms ([12](#)) converge to the system of differential equations ([6](#)) in [Theorem 1](#).

Given any pure strategy profile, the Gaussian belief updating process ([4](#)) implies that the agents' beliefs after every length- t history can be summarized by

$$(\mu(t), (\hat{\mu}_1(t), \dots, \hat{\mu}_N(t))),$$

where $\mu(t)$ and $\hat{\mu}_i(t)$ are the mean of the public belief and agent i 's private belief at time t , respectively. These beliefs are constant over intervals $[k\Delta, (k+1)\Delta)$, $k = 0, \dots, K-1 = (T-\Delta)/\Delta$. Moreover, over such intervals, the marginal contribution of each agent's effort

²⁷See [Cetemen \(2018\)](#) for an analysis of dynamic team behavior when members have different information about the project quality.

to feedback $Y((k+1)\Delta)$ and to final output P are constant. Since, in addition, effort cost is convex, we conclude that in any equilibrium, the effort choices are constant over such intervals. Thus, the setup is equivalent to a discrete time game with $K \equiv T/\Delta$ periods. Accordingly, we refer to the time interval $[\Delta k, \Delta(k+1))$ as period k . For ease of notation, let μ_k , $\hat{\mu}_{ik}$, and a_{ik} ($k = 0, \dots, K-1$) stand for the values of $\mu(t)$, $\hat{\mu}_i(t)$, and $a_i(t)$ for $t \in [k\Delta, (k+1)\Delta)$.

It is convenient to define

$$\tilde{Z}_k = \frac{Z((k+1)\Delta) - Z(k\Delta)}{\Delta} \quad \text{and} \quad \rho_k = \frac{\eta\Delta}{\nu_k}.$$

Then the evolution of the public belief can be expressed recursively as

$$\mu_{k+1} = (1 - \rho_{k+1})\mu_k + \rho_{k+1}\tilde{Z}_k. \quad (11)$$

The following proposition characterizes the unique equilibrium when feedback is observed at discrete intervals.

PROPOSITION 9. *Fix $\Delta > 0$. There exists a unique Nash equilibrium of the model. In equilibrium, agent i 's effort level for $t \in [k\Delta, (k+1)\Delta)$ ($k = 0, \dots, K-1$) is*

$$a_{ik}^* = \xi_{ik}\hat{\mu}_{ik},$$

where $\xi_{i,K-1} = s_i/c_i$, and

$$\xi_{ik} = \frac{s_i}{c_i} \left[1 + \kappa \sum_{l=k+1}^{K-1} \sum_{j \neq i} \xi_{jl} \rho_l \prod_{m=k+1}^{l-1} (1 - \kappa \xi_{im} \rho_m) \right] \quad (12)$$

for $k = 0, \dots, K-1$.

PROOF. We employ backward induction to prove the proposition. In the last period ($k = K-1$), each agent solves the problem

$$\begin{aligned} a_{i,K-1}^* &= \arg \max_a \mathbb{E} \left[\Delta s_i \theta \left(a + \sum_{j \neq i} a_{j,K-1}^* \right) \right] - \Delta c_i \frac{a^2}{2} \\ &= \arg \max_a s_i \hat{\mu}_{i,K-1} \left(a + \sum_{j \neq i} a_{j,K-1}^* \right) - c_i \frac{a^2}{2}. \end{aligned}$$

The first-order condition yields agent i 's unique equilibrium effort $a_{i,K-1}^* = (s_i/c_i)\hat{\mu}_{i,K-1}$, which is linear in the mean of the private belief with coefficient $\xi_{i,K-1} = s_i/c_i$.

Now suppose that the claim of the proposition holds for period $k+1$ onward, that is, in an equilibrium of the game, agent i plays $a_{il}^* = \xi_{il}\hat{\mu}_{il}$ for $l = k+1, \dots, K-1$, where ξ_{il} is defined in (12).

Fix a public history of length Δk , and suppose that there exists an equilibrium in which $\bar{a}_{ik}$ is agent i 's equilibrium effort choice in period k following the public history,

provided that he has not deviated in the past. Thus, all other agents anticipate that agent i chooses $\bar{a}_{ik}$. Then it suffices to show that $\bar{a}_{ik} = \xi_{ik} \hat{\mu}_{ik}$.

We compute agent i 's payoff from choosing arbitrary a in period k . First, using (11) we have

$$\begin{aligned} \mu_{k+1} - \hat{\mu}_{i,k+1} &= (1 - \rho_{k+1})(\mu_k - \hat{\mu}_{ik}) + \kappa \rho_{k+1}(a - \bar{a}_{ik}) \\ &= \frac{\rho_{k+1}}{\rho_k}(\mu_k - \hat{\mu}_{ik}) + \kappa \rho_{k+1}(a - \bar{a}_{ik}). \end{aligned} \quad (13)$$

To calculate the belief divergence in period $k + 2$ onward, we use the induction hypothesis that in period $l = k + 1, \dots, K - 1$ agent i plays $a_{il} = \xi_{il} \hat{\mu}_{il}$, while the others expect him to play $\bar{a}_{il} = \xi_{il} \mu_l$. Then, for all $l = k + 1, \dots, K - 1$,

$$\mu_{l+1} - \hat{\mu}_{i,l+1} = (1 - \rho_{l+1})(\mu_l - \hat{\mu}_{il}) - \kappa \rho_{l+1} \xi_{il}(\mu_l - \hat{\mu}_{i,l}) = \frac{\rho_{l+1}}{\rho_l} [1 - \kappa \rho_l \xi_{il}](\mu_l - \hat{\mu}_{il}),$$

from which we obtain

$$\begin{aligned} \mu_{l+1} - \hat{\mu}_{i,l+1} &= (\mu_{k+1} - \hat{\mu}_{i,k+1}) \prod_{m=k+1}^l \frac{\rho_{m+1}}{\rho_m} [1 - \kappa \rho_m \xi_m] \\ &= (\mu_{k+1} - \hat{\mu}_{i,k+1}) \frac{\rho_{l+1}}{\rho_{k+1}} \prod_{m=k+1}^l [1 - \kappa \rho_m \xi_m]. \end{aligned} \quad (14)$$

Substituting (13) into (14) and rearranging, we obtain

$$\mu_{l+1} = \hat{\mu}_{i,l+1} + \left(\frac{1}{\kappa \rho_k} (\mu_k - \hat{\mu}_{ik}) + (a - \bar{a}_{ik}) \right) \rho_{l+1} \kappa \prod_{m=k+1}^l [1 - \kappa \rho_m \xi_m]. \quad (15)$$

Now agent i 's optimal effort a^* solves

$$\begin{aligned} a^* &= \arg \max_a s_i \hat{\mu}_{ik} \left(a + \sum_{j \neq i} a_{jk}^* \right) - c_i \frac{a^2}{2} \\ &\quad + \mathbb{E}_k \left[s_i \sum_{l=k+1}^{K-1} \hat{\mu}_{il} \left(\xi_{il} \hat{\mu}_{il} + \sum_{j \neq i} \xi_{jl} \mu_l \right) - c_i \frac{\xi_{il}^2}{2} \hat{\mu}_{il}^2 \right]. \end{aligned}$$

Note that $\hat{\mu}_{il}$ for $l = k + 1, \dots, K - 1$ is independent of a and has expectation $\hat{\mu}_{ik}$. Substituting μ_l with (13) and (15), eliminating additive terms that are independent of a , and replacing $\hat{\mu}_{il}$ with its expectation whenever appropriate, agent i 's problem can be rewritten as

$$a^* = \arg \max_a s_i \hat{\mu}_{ik} \left[1 + \kappa \sum_{l=k+1}^{K-1} \sum_{j \neq i} \xi_{jl} \rho_l \prod_{m=k+1}^{l-1} (1 - \kappa \xi_{im} \rho_m) \right] a - c_i \frac{a^2}{2}.$$

It is clear that the problem is concave in a and has a unique solution. The first-order condition immediately yields the desired result.

Now define $\xi_i^\Delta(\cdot)$ as a step function such that for $t \in [k\Delta, (k+1)\Delta)$,

$$\xi_i^\Delta(t) = \xi_k.$$

Then, by [Proposition 9](#), $\xi_i^\Delta(\cdot)$ constitutes the unique equilibrium of the game where the feedback is observed at intervals of length Δ . It remains to show that as $\Delta \rightarrow 0$, the limit of $\xi_i^\Delta(\cdot)$ satisfies (6). Rewriting (12) in the recursive form yields

$$\xi_{ik} = \frac{s_i}{c_i} + \left[\kappa \frac{s_i}{c_i} \sum_{j \neq i} \xi_{j,k+1} \rho_{k+1} + (1 - \kappa \xi_{i,k+1} \rho_{k+1}) \left(\xi_{i,k+1} - \frac{s_i}{c_i} \right) \right].$$

Rearranging and substituting for ρ_k and $\xi_i^\Delta(\cdot)$, we obtain

$$\xi_i^\Delta(t) - \xi_i^\Delta(t + \Delta) = \kappa \rho(t + \Delta) \Delta \left[\frac{s_i}{c_i} \sum_{j \neq i} \xi_j^\Delta(t + \Delta) - \xi_i^\Delta(t + \Delta) \left(\xi_i^\Delta(t + \Delta) - \frac{s_i}{c_i} \right) \right].$$

Then dividing by Δ and letting $\Delta \rightarrow 0$, we obtain the desired result. $\square$

Proof of Proposition 1

To establish monotonicity, we first observe that $\dot{\xi}_i(T) < 0$. Suppose, for a contradiction, that there exist i and $\tilde{t} \in [0, T)$ such that $\dot{\xi}_i(\tilde{t}) > 0$. By the continuity of $\dot{\xi}_i(t)$, there exists $\hat{\Delta}_i > 0$ such that $\dot{\xi}_i(t) < 0$ for $t \in (T - \hat{\Delta}_i, T]$ and $\dot{\xi}_i(T - \hat{\Delta}_i) = 0$. Without loss of generality, assume that $i = 1$ attains $\min\{\hat{\Delta}_i | i = 1, \dots, N\}$, with the convention that $\hat{\Delta}_j = \infty$ if $\dot{\xi}_j(t) < 0$ for all t . This in particular implies that $\dot{\xi}_j(T - \hat{\Delta}_1) \leq 0$ and $\dot{\xi}_i(t) < 0$ for $t > T - \hat{\Delta}_1$ for all $j \neq 1$.

Step 1. Suppose that for all $j = 1, \dots, N$, $\dot{\xi}_j(T - \hat{\Delta}_1) = 0$. In this case, manipulation of (6) reveals that each $\xi_j(T - \hat{\Delta}_1)$ must be equal to its upper bound given in (1). Since they are all nonincreasing over the interval $(T - \hat{\Delta}_1, T]$, they must indeed be constant. In particular, it must be true that $\xi_j(T)$ is equal to its upper bound, which violates the terminal condition $\xi_j(T) = s_j/c_j$. Therefore, we conclude that there exists $j \neq 1$ such that $\dot{\xi}_j(T - \hat{\Delta}_1) < 0$.

Step 2. Next we claim that $\ddot{\xi}_1(T - \hat{\Delta}_1) > 0$. By taking derivatives of both sides of (6), and using $\dot{\xi}_1(T - \hat{\Delta}_1) = 0$ and (6), we obtain

$$\ddot{\xi}_1(T - \hat{\Delta}_1) = -\frac{\eta\kappa}{\nu_0 + \eta t} \frac{s_1}{c_1} \sum_{j=1}^N \dot{\xi}_j(T - \hat{\Delta}_1).$$

Then, since $\dot{\xi}_j(T - \hat{\Delta}_1) \leq 0$ with at least one strict inequality, we conclude that $\ddot{\xi}_1(T - \hat{\Delta}_1) > 0$. Now, since $\dot{\xi}_1$ is continuous, $\dot{\xi}_1(T - \hat{\Delta}_1) = 0$, and $\ddot{\xi}_1(T - \hat{\Delta}_1) > 0$, there exists $\epsilon > 0$ such that $\dot{\xi}_1(t) > 0$ whenever $t \in (T - \hat{\Delta}_1, T - \hat{\Delta}_1 + \epsilon)$, a contradiction, establishing that $\dot{\xi}_t \leq 0$ for all t .

Monotonicity immediately implies the lower bound on $\xi_i(t)$. Again, $\dot{\xi}_i(t) \leq 0$ implies, by (6), that the term in parentheses on the right-hand side must be nonnegative.

That is,

$$\frac{s_i}{c_i} \sum_{j=1}^N \xi_j(t) \geq \xi_i(t)^2.$$

Taking the square roots of both sides and summing over i yields

$$\sqrt{\sum_{j=1}^N \xi_j(t)} \leq \sum_{j=1}^N \sqrt{\frac{s_j}{c_j}}.$$

Combining the above two inequalities yields the desired upper bound.

The above argument implies that a solution remains bounded in finite time and, therefore, a solution exists for all $T \in [0, \infty)$. Next we show that there exists a unique solution to (6), which defines an autonomous first-order nonlinear system of differential equations. Define $\mathbf{u}(t) = (\xi_1(t), \dots, \xi_N(t), \nu(t))$. Then $d\mathbf{u}(t) = F(\mathbf{u}(t))$, where F is a Lipschitz continuous function given any (κ, η) . Then, by the Picard–Lindelöf theorem (Teschl 2012, Theorem 2.2), there exists a unique solution to this system in the domain $[0, T]$ with the boundary values $\xi_i(T) = s_i/c_i$ for all i . $\square$

Proof of Propositions 2 and 3

Note that $\rho(t) = \eta/(\nu_0 + \eta t)$ uniformly increases in η and decreases in ν_0 . Next, take two appropriate functions $\rho^a(\cdot)$ and $\rho^b(\cdot)$ such that for all $t \in [0, T]$, $\rho^a(t) < \rho^b(t)$. Let $\{\xi_i^a(t)\}_{i=1}^N$ and $\{\xi_i^b(t)\}_{i=1}^N$ be the solutions to (6) when $\rho(\cdot) = \rho^a(\cdot)$ and $\rho(\cdot) = \rho^b(\cdot)$, respectively. For all i , since $\xi_i^a(T) = \xi_i^b(T) = s_i/c_i$, then $\dot{\xi}_i^a(T) > \dot{\xi}_i^b(T)$, which implies that there exists $\varepsilon_i > 0$ ($i = 1, \dots, N$) such that $\xi_i^a(t) < \xi_i^b(t)$ for any $t \in (T - \varepsilon_i, T)$. Let $\varepsilon = \min_i \varepsilon_i$. Suppose there exist $\tilde{t}$ and i such that $\xi_i^a(\tilde{t}) \geq \xi_i^b(\tilde{t})$. Then by continuity of $\xi_i^a(t)$ and $\xi_i^b(t)$, there must exist j and $t^* \in [\tilde{t}, T - \varepsilon]$ such that $\xi_j^a(t^*) = \xi_j^b(t^*)$ and for all $t \in (t^*, T)$ and for all i , $\xi_i^a(t) < \xi_i^b(t)$. However, since for all i , $\xi_i^a(t^*) \leq \xi_i^b(t^*)$, $\xi_j^a(t^*) = \xi_j^b(t^*)$, and $\rho^a(t^*) < \rho^b(t^*)$, (6) implies that $\dot{\xi}_j^a(t^*) > \dot{\xi}_j^b(t^*)$. This in turn implies that there exists $\varepsilon' > 0$ such that $\xi_j^a(t) > \xi_j^b(t)$ for all $t \in (t^*, t^* + \varepsilon')$, a contradiction. This establishes that for all $t \in [0, T)$ and for all i , $\xi_i^a(t) < \xi_i^b(t)$. Combined with the monotonicity of $\rho(\cdot)$ in η and ν_0 , this completes the proof of Proposition 2 and item (ii) of Proposition 3.

The proof that $\xi_i(t)$ increases in κ is analogous. For the convergence result in item (i), note that by (6), for all i , $\dot{\xi}_i(T) = -\rho(T)\kappa(s_i/c_i) \sum_{j \neq i} (s_j/c_j)$, which approaches $-\infty$ as $\kappa \rightarrow \infty$. Together with strict monotonicity of $\xi_i(\cdot)$, this establishes the pointwise convergence to the upper bound.

For item (iii), fix two constants $\tilde{c}_i$ and $\hat{c}_i$ with $0 < \tilde{c}_i < \hat{c}_i$, and let $\{\tilde{\xi}_i(t)\}_{i=1}^N$ and $\{\hat{\xi}_i(t)\}_{i=1}^N$ be the solutions to (6) when $c_i = \tilde{c}_i$ and $c_i = \hat{c}_i$, respectively. Then it suffices to show that $\tilde{\xi}_j(t) > \hat{\xi}_j(t)$ for all $t \in [0, T)$ and $j = 1, \dots, N$.

Take an increasing sequence c_i^k with $c_i^0 = \tilde{c}_i$ and $\lim_{k \rightarrow \infty} c_i^k = \hat{c}_i$. Let $J = \{1, \dots, i-1, i+1, \dots, N\}$. For $j \in J \cup \{i\}$, define $\xi_j^0 \equiv \tilde{\xi}_j$. Then, for $k \geq 1$, define ξ_j^k recursively as follows:

- The ξ_i^k solves (6) for $c_i = c_i^k$ while keeping $\xi_j(\cdot)$ fixed at $\xi_j^{k-1}(\cdot)$ for all $j \in J$. Formally, ξ_i^k satisfies

$$\dot{\xi}_i^k(t) = -\rho(t)\kappa \left[\frac{s_i}{c_i^k} \sum_{j \neq i} \xi_j^{k-1}(t) - \xi_i(t) \left(\xi_i(t) - \frac{s_i}{c_i^k} \right) \right], \quad (16)$$

with a terminal condition $\xi_i^k(T) = s_i/c_i^k$. That is, ξ_i^k is agent i 's best response to $\{\xi_j^{k-1}\}_{j \in J}$ when his marginal effort cost is c_i^k .

- Given $\xi_i^k(\cdot)$, $\{\xi_j^k\}_{j \in J}$ solves (6), assuming that $\xi_i(\cdot) = \xi_i^k(\cdot)$. Formally, $\{\xi_j^k\}_{j \in J}$ satisfies the system of $N - 1$ differential equations

$$\dot{\xi}_j^k(t) = -\rho(t)\kappa \left[\frac{s_j}{c_j} \left(\xi_i^k(t) + \sum_{l \neq j} \xi_l(t) \right) - \xi_j(t) \left(\xi_j(t) - \frac{s_j}{c_j} \right) \right], \quad (17)$$

with terminal conditions $\xi_j^k(T) = s_j/c_j$. Note that $\{\xi_j^k\}_{j \in J}$ is the equilibrium strategy profile in a game played among agents in J , with agent i 's strategy fixed at ξ_i^k .

Using an inductive argument, we establish that for all $k \geq 0$,

$$\begin{aligned} \xi_i^{k+1}(t) &< \xi_i^k(t) \quad \text{for } t \in [0, T], \\ \xi_j^{k+1}(t) &< \xi_j^k(t) \quad \text{for } t \in [0, T] \text{ and } j \in J. \end{aligned}$$

We conduct our analysis in the following three steps.

STEP 1. For all $t \in [0, T]$, $\xi_i^1(t) < \xi_i^0(t)$.

PROOF. Note that $\xi_i^1(T) < \xi_i^0(T)$ from the boundary condition. Suppose to the contrary that there exists $\tilde{t} \in [0, T)$ such that $\xi_i^1(\tilde{t}) \geq \xi_i^0(\tilde{t})$. Then since $\xi_i^k(t)$ is continuous in t , there must exist $t^* \in [\tilde{t}, T)$ such that $\xi_i^1(t^*) = \xi_i^0(t^*)$ and $\xi_i^1(t) < \xi_i^0(t)$ for all $t \in (t^*, T]$. However, since $c_i^1 > c_i^0$, from (16) we have $\dot{\xi}_i^1(t^*) > \dot{\xi}_i^0(t^*)$, which in turn implies that there exists $\varepsilon > 0$ such that $\xi_i^1(t) > \xi_i^0(t)$ for all $t \in (t^*, t^* + \varepsilon)$, leading to a contradiction. $\square$

STEP 2. For all $t \in [0, T)$ and $j \in J$, $\xi_j^1(t) < \xi_j^0(t)$.

PROOF. From the boundary conditions, $\xi_j^1(T) = \xi_j^0(T)$ for all $j \in J$. Since $\xi_i^1(T) > \xi_i^0(T)$, from (17) we have $\dot{\xi}_j^1(T) < \dot{\xi}_j^0(T)$, implying that there exists $\varepsilon > 0$ such that for all $j \in J$ and $t \in (T - \varepsilon, T)$, $\xi_j^1(t) > \xi_j^0(t)$. Suppose to the contrary that there exist $j \in J$ and $\tilde{t} \in [0, T - \varepsilon]$ such that $\xi_j^1(\tilde{t}) \geq \xi_j^0(\tilde{t})$. Then, by continuity of $\xi_j^k(\cdot)$, there must exist $l \in J$ and $t^* \in [\tilde{t}, T - \varepsilon)$ such that $\xi_l^1(t^*) = \xi_l^0(t^*)$ and $\xi_j^1(t) < \xi_j^0(t)$ for all $j \in J$ and $t \in (t^*, T)$. However, since $\xi_i^1(t^*) > \xi_i^0(t^*)$ (by Step 1) and $\xi_j^1(t^*) \geq \xi_j^0(t^*)$ for all $j \in J$ (by continuity of $\xi_j^k(\cdot)$), from (17) we have $\dot{\xi}_l^1(t^*) > \dot{\xi}_l^0(t^*)$, which in turn implies that there exists $\varepsilon' > 0$ such that $\xi_l^1(t) > \xi_l^0(t)$ for all $t \in (t^*, t^* + \varepsilon')$, leading to a contradiction. $\square$

STEP 3 (Induction). Suppose that for some $\bar{k} \in \mathbb{N}$, $\xi_i^{\bar{k}}(t) < \xi_i^{\bar{k}-1}(t)$ for $t \in [0, T]$ and $\xi_j^{\bar{k}}(t) < \xi_j^{\bar{k}-1}(t)$ for $t \in [0, T]$ and $j \in J$. Then $\xi_i^{\bar{k}+1}(t) < \xi_i^{\bar{k}}(t)$ for $t \in [0, T]$ and $\xi_j^{\bar{k}+1}(t) < \xi_j^{\bar{k}}(t)$ for $t \in [0, T]$ and $j \in J$.

PROOF. The proof is almost identical to those in Steps 1 and 2. First consider $\xi_i^{\bar{k}+1}$ versus $\xi_i^{\bar{k}}$. Note that $\xi_i^{\bar{k}+1}(T) < \xi_i^{\bar{k}}(T)$ from the boundary condition. Suppose to the contrary that there exists $\tilde{t} \in [0, T)$ such that $\xi_i^{\bar{k}+1}(\tilde{t}) \geq \xi_i^{\bar{k}}(\tilde{t})$. Then since $\xi_i^{\bar{k}}(t)$ is continuous in t , there must exist $t^* \in [\tilde{t}, T)$ such that $\xi_i^{\bar{k}+1}(t^*) = \xi_i^{\bar{k}}(t^*)$ and $\xi_i^{\bar{k}+1}(t) < \xi_i^{\bar{k}}(t)$ for all $t \in (t^*, T]$. However, since $c_i^1 > c_i^0$ and $\xi_j^{\bar{k}}(t) < \xi_j^{\bar{k}-1}(t)$ for all $j \in J$, from (16) we have $\dot{\xi}_i^{\bar{k}+1}(t^*) > \dot{\xi}_i^{\bar{k}}(t^*)$, which in turn implies that there exists $\varepsilon > 0$ such that $\xi_i^{\bar{k}+1}(t) > \xi_i^{\bar{k}}(t)$ for all $t \in (t^*, t^* + \varepsilon)$, a contradiction. Finally, an argument identical to Step 2 proves that $\xi_j^{\bar{k}+1}(t) < \xi_j^{\bar{k}}(t)$ for all $j \in J$ and $t \in [0, T]$. $\square$

Given the arguments in Steps 1–3, we finish the proof by showing that for any $j \in J \cup \{i\}$ and $t \in [0, T]$,

$$\lim_{k \rightarrow \infty} \xi_j^k(t) = \hat{\xi}_j(t).$$

For any t and j , $\{\xi_j^k(t)\}_{k=0}^\infty$ is a decreasing and bounded sequence, and, therefore, it converges. Then the fact that the limit coincides with $\hat{\xi}_j(t)$ follows by continuity of the system (6).

For the last part we assume agents are symmetric and each has the same costs and shares; then we can solve $\xi^N(t)$ in closed form as

$$\xi^N(t) = \frac{1}{c + c(N-1)(c(\eta_0 + \eta t))^{\frac{\kappa}{c}}(c(\eta_0 + \eta T))^{\frac{-\kappa}{c}}}$$

or, alternatively,

$$\xi^N(t) = \frac{1}{c + c(N-1)(c\rho(t))^{\frac{\kappa}{c}}(c\rho(T))^{\frac{-\kappa}{c}}}.$$

Then as $N \rightarrow \infty$, $\xi^N(t) \rightarrow 0$ pointwise. However, the aggregate effort coefficient $(\xi_t^N N)$ converges pointwise to

$$\lim_{N \rightarrow \infty} N\xi^N(t) = \left(\frac{\rho(T)}{\rho(t)} \right)^{\frac{\kappa}{c}}.$$

Note that the term inside the parentheses is always greater than 1 and decreasing at t . In addition, if we examine the derivative of the term $N\xi^N(t)$ with respect to N , we obtain

$$\frac{-(c\rho(t))^{\frac{\kappa}{c}}(c\rho(T))^{\frac{\kappa}{c}} + (c\rho(T))^{\frac{2\kappa}{c}}}{c((c(N-1)\rho(t))^{\frac{\kappa}{c}} + (c\rho(T))^{\frac{\kappa}{c}})^2},$$

which is positive all the time.

Proof of Proposition 5

Recall from (10) that the expected output under the feedback schedule $\mathbb{T}$ is

$$P(\mathbb{T}) = \int_0^T \left(\mu_0^2 + \frac{\eta \tau_{\mathbb{T}}(t)}{\nu_0(\nu_0 + \eta \tau_{\mathbb{T}}(t))} \right) \sum_{i=1}^N \xi_i^{\mathbb{T}}(t) dt,$$

where $\tau_{\mathbb{T}}(t) = \sup\{t' \in \mathbb{T} | t' \leq t\}$ is the most recent date of feedback preceding time t . Consider a one-time feedback schedule $\mathbb{T} = \{\hat{t}\}$ for some $\hat{t} \in (0, T)$. Then from (12) in Proposition 9, agent i 's belief sensitivity of effort is given by

$$\xi_i^{\mathbb{T}}(\hat{t}) = \begin{cases} \frac{s_i}{c_i} \left(1 + \frac{\kappa \eta (T - \hat{t})}{\nu_0 + \eta \hat{t}} \sum_{j \neq i} \frac{s_j}{c_j} \right) & \text{for } t < \hat{t}, \\ \frac{s_i}{c_i} & \text{for } t \geq \hat{t}. \end{cases}$$

Therefore, the expected output is

$$P(\{\hat{t}\}) = \hat{t} \mu_0^2 \left(\sum_i \frac{s_i}{c_i} + \frac{2\kappa \eta (T - \hat{t})}{\nu_0 + \eta \hat{t}} \sum_{i \neq j} \frac{s_i s_j}{c_i c_j} \right) + (T - \hat{t}) \left(\mu_0^2 + \frac{\eta \hat{t}}{\nu_0(\nu_0 + \eta \hat{t})} \right) \sum_i \frac{s_i}{c_i}. \quad (18)$$

Let $t^* = \arg \max_{\hat{t}} P(\{\hat{t}\})$. Then by the first-order condition, t^* must solve

$$\frac{-\eta}{(\nu_0 + \eta t^*)^2} \left(2\mu_0^2 \kappa \sum_{i \neq j} \frac{s_i s_j}{c_i c_j} + \frac{1}{\nu_0} \sum_i \frac{s_i}{c_i} \right) (\eta (t^*)^2 + 2\nu_0 t^* - \nu_0 T) = 0,$$

which has a unique positive solution

$$t^* = \frac{-\nu_0 + \sqrt{\nu_0^2 + \eta \nu_0 T}}{\eta}.$$

Let $\mathbb{T}^{**} = [0, T]$ be the continuous feedback schedule. Then

$$P(\mathbb{T}^{**}) = \int_0^T \left(\mu_0^2 + \frac{\eta t}{\nu_0(\nu_0 + \eta t)} \right) \sum_{i=1}^N \xi_i(t) dt,$$

where $\xi_i(t)$ solves (6) in Theorem 1. By Proposition 1, $\xi_i(t)$ is bounded above for any $t \in [0, T]$, implying that $P(\mathbb{T}^{**})$ is bounded. However, (18) implies that $P(\{t^*\})$ diverges to infinity as $\kappa \rightarrow \infty$ and, thus, $P(\{t^*\}) > P(\mathbb{T}^{**})$ for sufficiently large κ . $\square$

Proof of Proposition 6

First, we write (12) in a recursive form:

$$\xi_{i,k} = \xi_{i,k+1} + \kappa \rho_{k+1} \frac{s_i}{c_i} \sum_{j \neq i} \xi_{j,k+1} - \kappa \rho_{k+1} \xi_{i,k+1}^2. \quad (19)$$

Fix any feedback schedule $\mathbb{T} = \{t_1, \dots, t_l\}$ with $l \geq 2$ and $t_{m+1} - t_m \geq \Delta$ for any $m = 1, \dots, l-1$. Note that having l instances of feedback divides $[0, T]$ into $(l+1)$ *periods*. Then the above recursion implies that there exists $\bar{\kappa}$ such that for any $\kappa \geq \bar{\kappa}$ and i , $\xi_{i,l-2} < 0$. In addition, it is straightforward from (19) that if $\xi_{i,m} < 0$ for all i and for some m , then $\xi_{i,k} < 0$ for all i and $k = 1, \dots, m-1$. Moreover, $\xi_{i,k-1}$ is always in a higher order than $\xi_{i,k}$ for all $k \in \{1, \dots, l\}$. Let $t_0 = 0$ and $t_{l+1} = T$. Then from (10), the expected total output under $\mathbb{T}$ is given by

$$P(\mathbb{T}) = \sum_{m=0}^l (t_{m+1} - t_m) \left(\mu_0^2 + \frac{\eta t_m}{\nu_0(\nu_0 + \eta t_m)} \right) \sum_{i=1}^N \xi_{i,m}.$$

Now consider the one-time feedback schedule $\hat{\mathbb{T}} = \{t_l\}$. We claim that the expected output under $\hat{\mathbb{T}}$ is greater than that under $\mathbb{T}$ for sufficiently large κ . Observe that total output for $t \in [t_l, T]$ does not change. If we switch the feedback schedule from $\mathbb{T}$ to $\hat{\mathbb{T}}$, the loss from having less precise information for $t \in [t_{l-1}, t_l]$ is

$$(t_l - t_{l-1}) \frac{\eta t_{l-1}}{\nu_0(\nu_0 + \eta t_{l-1})} \sum_{i=1}^N \xi_{i,l-1}.$$

However, for any $t < t_{l-1}$, $\xi_i(t)$ is negative under $\mathbb{T}$ but is positive under $\hat{\mathbb{T}}$. Therefore, from (19), the gain from higher $\xi_i(t)$ is at least

$$-t_{l-1} \kappa \rho_{l-1} \sum_{i=1}^N \left(\frac{s_i}{c_i} \sum_{j \neq i} \xi_{j,l-1} - \xi_{i,l-1}^2 \right).$$

Since the loss is of order κ and the gain is at least of order κ^3 , the expected output under $\hat{\mathbb{T}}$ is greater than that under $\mathbb{T}$ for sufficiently large κ . $\square$

Proof of Proposition 8

From (10), the expected output of a team given performance measure scheme $\zeta = \text{IPM, JPM}$ and feedback schedule $\mathbb{T}$ is given by

$$P(\zeta, \mathbb{T}) = \int_0^T \left(\mu_0^2 + \frac{\eta \tau_{\mathbb{T}}(t)}{\nu_0(\nu_0 + \eta \tau_{\mathbb{T}}(t))} \right) \sum_{i=1}^N \xi_i^{\zeta, \mathbb{T}}(t) dt,$$

where $\tau_{\mathbb{T}}(t) = \sup\{t' \in \mathbb{T} | t' \leq t\}$ is the most recent date of feedback preceding time t .

For the scheme $(\text{IPM}, \mathbb{T}^{**})$, it is straightforward that $\xi_i^{\text{IPM}, \mathbb{T}^{**}}(t) = 1/c_i$ for all $t \in [0, T]$ and $\tau_{\mathbb{T}^{**}}(t) = t$. Therefore,

$$P(\text{IPM}, \mathbb{T}^{**}) = \int_0^T \left(\mu_0^2 + \frac{\eta t}{\nu_0(\nu_0 + \eta t)} \right) dt \cdot \sum_{i=1}^N \frac{1}{c_i}.$$

From the argument in the proof of [Proposition 5](#), the expected output under the scheme (JPM, $\mathbb{T}^*$) is

$$P(\text{JPM}, \mathbb{T}^*) = t^* \mu_0^2 \sum_{i=1}^N \hat{\xi}_i + (T - t^*) \left(\mu_0^2 + \frac{\eta t^*}{\nu_0(\nu_0 + \eta t^*)} \right) \sum_{i=1}^N \frac{s_i}{c_i},$$

where $\hat{\xi}_i$ s are given by

$$\hat{\xi}_i = \frac{s_i}{c_i} \left(1 + \frac{\kappa \eta (T - t^*)}{\nu_0 + \eta t^*} \sum_{j \neq i} \frac{s_j}{c_j} \right).$$

Note that as $\kappa \rightarrow \infty$, $\hat{\xi}_i \rightarrow \infty$ for all i . Therefore, $P(\text{JPM}, \mathbb{T}^*) > P(\text{IPM}, \mathbb{T}^{**})$ for sufficiently large κ , showing the desired result. $\square$

APPENDIX B: CONTINUOUS-TIME MODEL AND EQUILIBRIUM CONVERGENCE

In this appendix, we set up a full continuous-time model in which agents observe feedback at every instant rather than at the discrete-time intervals of $\Delta > 0$. In this model, agents react to information at every instant and change their actions accordingly. We characterize the equilibrium behavior in this model and show that this coincides with the unique Nash equilibrium in the continuous-feedback limit ($\Delta \rightarrow 0$) of our main model.

We focus on linear Markovian strategies and construct an equilibrium that is unique in the class of linear Markov perfect equilibria.²⁸ Formally, we assume that each agent uses a linear strategy of the form $a^*(\mu) = \xi_{i,t} \mu_t$. Define $\delta_i(t) = \mu(t) - \hat{\mu}_i(t)$ as the difference between public and private belief from agent i 's perspective. Assume that each agent except agent i follows the equilibrium strategy $a^*(\mu)$.

Similar to the main section, we use the private belief and the belief difference between the public and the private belief as state variables for the agent's problem. The evolution of the agent's private belief is a standard filtering problem, which enables us to apply the method of [Liptser and Shiryaev \(2013\)](#) (Theorem 12.1). Given the evolution of belief and the conjectured equilibrium strategy for the other agents, agent i 's best-response problem becomes a standard stochastic optimal control problem. Therefore, agent i 's problem can be described recursively by the Hamilton–Jacobi–Bellman (HJB) equation

$$\begin{aligned} 0 = \sup_{a \in \mathbb{R}} & \left\{ s_i \hat{\mu}_i \left(a + \sum_{j \neq i} \xi_j(t) (\delta_j(t) + \hat{\mu}_j(t)) \right) - c_i \frac{1}{2} a^2 + \frac{1}{2} \frac{\eta^2}{\nu_t^2} V_{\hat{\mu}, \hat{\mu}}(\hat{\mu}, \delta) \right. \\ & \left. + \left[\frac{\eta}{\nu_t} \kappa (a - \xi_i(\delta_i(t) + \hat{\mu}_i(t))) - \frac{\eta}{\nu_t} \delta \right] V_{\delta}(\hat{\mu}, \delta) + V_t \right\}, \\ d\hat{\mu}_i(t) &= \frac{\eta}{\nu(t)} dZ^i(t), \end{aligned}$$

²⁸We conjecture that it is the unique equilibrium of the game with continuous feedback, although we do not have a formal proof.

$$\dot{\delta}_i(t) = \left(-\frac{\eta}{\nu_t} \delta_t + \frac{\eta}{\nu_t} \kappa (a_t - \xi_i(t)(\delta_i(t) + \hat{\mu}_i(t))) \right) dt,$$

$$\dot{\nu}(t) = \eta,$$

where $Z^i(t) = \sqrt{\eta}(Y(t) - \kappa \int_0^t (a_i(t') + \sum_{j \neq i} a_j^*(t')) dt')$. Since the flow payoff is quadratic, it is natural to guess a quadratic value function of the form $v_{1,i}(t)\hat{\mu}^2 + v_{2,i}(t)\hat{\mu}\delta + v_{3,i}(t)$ for the HJB equation. Then by the first-order condition with respect to $a_i(t)$, we have the equality

$$a_i(t) = \frac{s_i}{c_i} \hat{\mu} + \frac{\eta \kappa}{\nu(t)c_i} v_{2,i}(t) \hat{\mu} \Rightarrow v_{2,i}(t) = \frac{\left(\xi_i(t) - \frac{s_i}{c_i} \right) \nu(t) c_i}{\eta \kappa}.$$

Taking the time derivative of both sides and rearranging gives

$$\dot{v}_{2,i}(t) = \frac{1}{\eta \kappa} (\dot{\xi}_i(t) \nu(t) c_i + \dot{\nu}(t) \xi_i(t) c_i - \dot{\nu}(t) s_i) = \frac{\dot{\xi}_i(t) \nu(t) c_i}{\eta \kappa} + \xi_i(t) \frac{1}{\kappa} c_i - s_i \frac{1}{\kappa}.$$

Then applying the envelope theorem to the HJB equation (evaluated at $\delta_i = 0$) and plugging into the equations above, we reach²⁹

$$\frac{\left(\xi_i(t) - \frac{s_i}{c_i} \right) \nu(t) c_i}{\eta \kappa} \left(\frac{\eta}{\nu(t)} (\kappa \xi_i(t) + 1) \right) = s_i \sum_{j \neq i}^N \xi_j(t) + \frac{\dot{\xi}_i(t) \nu(t) c_i}{\eta \kappa} + \xi_i(t) \frac{1}{\kappa} c_i - s_i \frac{1}{\kappa}.$$

Rearranging yields

$$\dot{\xi}_i(t) = -\frac{\eta \kappa}{\nu_0 + \eta t} \left[\frac{s_i}{c_i} \sum_{j \neq i}^N \xi_j(t) - \xi_i(t) \left(\xi_i(t) - \frac{s_i}{c_i} \right) \right],$$

which is identical to the unique Nash equilibrium in the continuous-feedback limit of our main model.

The individual action $(a_i(t))$ follows an Ito process:

$$da_i(t) = \mu(t) \left(-\frac{\eta \kappa}{\nu(t)} \left(\frac{s_i}{c_i} \sum_j^N \xi_j(t) - (\xi_i(t))^2 \right) \right) dt + \frac{\eta}{\nu(t)} \xi_i(t) dZ^i(t).$$

Note that the drift is negative if and only if $\mu(t) > 0$, since the term within the parentheses on the right-hand side is always negative.³⁰ The volatility term $(\frac{\eta}{\nu(t)} \xi_i(t))$ is decreasing over time, since $\nu(t)$ is increasing, and $\xi(t)$ is decreasing. The volatility term converges to 0 as $T \rightarrow \infty$, given that $\xi(t) \rightarrow \frac{s_i}{c_i}$ and $\nu(t) \rightarrow \infty$. Therefore, if the time horizon is sufficiently long, the equilibrium converges to the equilibrium of the static game with complete information.

²⁹We evaluate the equation at $\delta_i = 0$ since the public belief coincides with the private belief on the equilibrium path.

³⁰Precisely, $Z^i(t)$ is a Brownian motion from the perspective of agent i .

Verification We can evaluate the parameters $v_{1,i}(t)$, $v_{2,i}(t)$, and $v_{3,i}(t)$ as follows. By matching the coefficients, we obtain

$$\begin{aligned}\dot{v}_{1,i}(t) &= -s_i \sum_j^N \xi_j(t) - c_i \frac{1}{2} \xi_i(t)^2, \\ \dot{v}_{2,i}(t) &= v_{2,i}(t) \left(\frac{\eta}{\nu_t} (\kappa \xi_i(t) + 1) \right) - s_i \sum_{j \neq i} \xi_j(t), \\ \dot{v}_{3,i}(t) &= -\frac{1}{2} \frac{\eta^2}{\nu_t^2} v_{1,i}(t).\end{aligned}$$

The parameter V is twice continuously differentiable in δ and $\hat{\mu}$, and continuously differentiable in t . Since the form is quadratic, it satisfies the polynomial growth condition. Then, by Theorem 3.1 (Fleming and Soner (2006), Chapter IV), we conclude that the conjectured form solves the agent's problem.

Convergence from the discrete-feedback model It remains to show that for all i , $a_i^\Delta(t)$ converges in distribution to $a_i(t)$. Recall that $a_i^\Delta(t) = \xi_i^\Delta(n) \mu_i^\Delta(n)$, where n is such that $t \in [n\Delta, (n+1)\Delta]$. The parameter $\xi_i^\Delta(t)$ converges pointwise to $\xi_i(t)$ and by the Donsker invariance principle, $\mu_i^\Delta(t)$ converges in distribution to $\mu_i^\Delta(t)$. Then by Whitt (1980) multiplication, $\xi_i^\Delta(t) \mu_i^\Delta(t)$ converges in distribution to $\xi_i(t) \mu_i(t)$.

APPENDIX C: NON-MONOTONICITY OF BELIEF SENSITIVITY

As discussed in the main text, when the period length Δ is away from 0, the belief sensitivity ξ_{it} may be non-monotonic over time. To see this, consider a special case with two symmetric agents (i.e., $c_1 = c_2 \equiv c/2$, and $s_1 = s_2 = 1/2$) and a horizon of three periods. Furthermore, let $\eta_0 = \eta = \kappa = 1$. Note that in this case, $\rho_t = \kappa/(t+2)$, $t = 0, 1, 2$. Letting $\xi_{1t} = \xi_{2t} \equiv \xi_t$, it is easy to calculate ξ_0 , ξ_1 , and ξ_2 as

$$\begin{aligned}\xi_2 &= \frac{1}{c}, \\ \xi_1 &= \frac{1}{c} \left[1 + \frac{\kappa}{1+3\Delta} \Delta \frac{1}{c} \right], \\ \xi_0 &= \frac{1}{c} \left[1 + \frac{\kappa}{1+2\Delta} \Delta \frac{1}{c} \left[1 + \frac{\kappa}{1+3\Delta} \Delta \frac{1}{c} \right] + \underbrace{\left(1 - \frac{\kappa}{1+2\Delta} \Delta \frac{1}{c} \left[1 + \frac{\kappa}{1+3\Delta} \Delta \frac{1}{c} \right] \right)}_{\text{last period return}} \frac{\kappa}{1+3\Delta} \Delta \frac{1}{c} \right].\end{aligned}$$

It is apparent by observation that ξ_1 is necessarily larger than ξ_2 , while ξ_0 can be larger or smaller than ξ_1 , depending on the values of Δ , κ , and c . In particular, for a fixed level of Δ , if κ is very large or c is very small, then ξ_0 falls below ξ_1 , since the term marked “last period return” becomes unboundedly negative. Intuitively, this is because in those cases, the ratchet effect in the middle period is large. To see this, notice that under the said conditions, ξ_1 can be very large. This means that the effort choice in the middle

period is very sensitive to beliefs, either because the marginal cost of effort is very small (i.e., small c) or the feedback is very sensitive to effort choice (i.e., large κ). This implies that the impact of a divergence between public and private beliefs on effort choices at the beginning of this period is greatly amplified in this period. In particular, after an upward deviation in period 0, player i in period 1, being less optimistic than his teammate, chooses an effort level that is far below the teammate's expectations. This in turn biases the period 2 belief of his teammate downward by a large amount and, thus, greatly reduces teammate's period 2 effort—in fact to a level below what it would have been without the period 0 deviation of player i . The impact of this reduction overcomes the impact of the increase in the teammate's effort in the middle period, rendering the marginal product of effort in the initial period small and possibly even negative.

It is worth noting that this non-monotonicity is an artifact of our discrete-time specification. In fact, when actions can be adjusted frequently and feedback is received frequently, this type of behavior disappears. This is easily observed by inspecting the expressions for ξ_0 , ξ_1 , and ξ_2 above. Intuitively, this is because the belief divergence after a deviation causes a ratchet effect that is the source of the non-monotonicity. The strength of the ratchet effect is positively related to the amount of information released in each period, which is proportional to the period length Δ .

APPENDIX D: ROLE OF PROJECT UNCERTAINTY

As discussed in [Section 4](#), as the project uncertainty increases, there exists a trade-off between benefit from the belief manipulation incentives and the standard cost due to uninformed effort choices. In this appendix, we analyze the effect of project uncertainty on the agents' expected payoff, which is used to plot [Figure 2](#).

Suppose that a manager of a team faces a choice of projects with varying uncertainty. The team manager tries to maximize the ex ante total payoff to the team. To clarify our analysis of the trade-off, we consider the case in which all projects have *the same ex ante value under complete information*. Recall that if the project state θ is perfectly observed at the beginning, then the equilibrium action is $a_i^*(t) = \theta/N$ for all $t \in [0, T]$. Since the state θ is normally distributed with mean μ_0 and precision ν_0 , the agent's ex ante expected payoff *before* the realization of θ is

$$\begin{aligned} \mathbb{E}_0 \left[\int_0^T \left(\theta \cdot a_i^*(t) - \frac{(a_i^*(t))^2}{2} \right) dt \right] &= \frac{T}{N} \left(1 - \frac{1}{2N} \right) \mathbb{E}_0[\theta^2] \\ &= \frac{T}{N} \left(1 - \frac{1}{2N} \right) \left(\mu_0^2 + \frac{1}{\nu_0} \right). \end{aligned}$$

Note that the payoff structure of our model implies that the value of the project is convex in θ . Therefore, choosing a risky project (one with a small ν_0) is always beneficial under complete information.

Now consider the original model where θ is unknown. We consider the optimal choice of uncertainty ν_0 subject to a constraint $\mu_0^2 + \frac{1}{\nu_0} = k$ for some $k > 0$. This constraint requires that the mean of the project decreases as its level of uncertainty increases, offsetting the inherent benefit of risk-taking described above. Given the constraint, the ex ante expected equilibrium payoff is given by

$$\begin{aligned}\mathbb{E}_0\left[\int_0^T\left(\theta \cdot a_i(t) - \frac{(a_i(t))^2}{2}\right)dt\right] &= \int_0^T \xi(t)\left(1 - \frac{\xi(t)}{2}\right)\mathbb{E}_0[\mu(t)^2]dt \\ &= \int_0^T \xi(t)\left(1 - \frac{\xi(t)}{2}\right)\left(k - \frac{1}{\nu(t)}\right)dt.\end{aligned}$$

Note that as the project uncertainty becomes larger, the cost of uncertainty (captured by the term $1/\nu(t)$) increases, while the free-riding problem is alleviated since $\xi(t)$ uniformly increases in ν_0 for all $t \in (0, T)$. Figure 2, which plots the above formula as a function of ν_0 , shows that there exists an optimal level of project uncertainty that balances the trade-off between the cost of uncertainty and the benefit of belief manipulation.

REFERENCES

- Admati, Anat R. and Motty Perry (1991), “Joint projects without commitment.” *Review of Economic Studies*, 58, 259–276. [1026]
- Alchian, Armen A. and Harold Demsetz (1972), “Production, information costs, and economic organization.” *American Economic Review*, 62, 777–795. [1026, 1035]
- Bandiera, Oriana, Iwan Barankay, and Imran Rasul (2013), “Team incentives: Evidence from a firm level experiment.” *Journal of the European Economic Association*, 11, 1079–1114. [1023]
- Bhaskar, Venkataraman (2014), “The ratchet effect re-examined: A learning perspective.” CEPR Discussion Paper 9956. [1025]
- Bhaskar, Venkataraman and George J. Mailath (2019), “The curse of long horizons.” *Journal of Mathematical Economics*, 82, 74–89. [1025]
- Bolton, Patrick and Christopher Harris (1999), “Strategic experimentation.” *Econometrica*, 67, 349–374. [1025, 1027, 1036]
- Bonatti, Alessandro and Johannes Hörner (2011), “Collaborating.” *American Economic Review*, 101, 632–663. [1026, 1029]
- Campbell, Arthur, Florian Ederer, and Johannes Spinnewijn (2014), “Delay and deadlines: Freeriding and information revelation in partnerships.” *American Economic Journal: Microeconomics*, 6, 163–204. [1027]
- Cetemen, Doruk (2018), “Achieving efficiency in repeated partnerships via information design.” Unpublished paper, Department of Economics, University of Rochester. [1043]
- Che, Yeon-Koo and Seung-Weon Yoo (2001), “Optimal incentives for teams.” *American Economic Review*, 91, 525–541. [1027, 1042]

Cisternas, Gonzalo (2018a), “Two-sided learning and the ratchet principle.” *Review of Economic Studies*, 85, 307–351. [1025, 1026]

Cisternas, Gonzalo (2018b), “Career concerns and the nature of skills.” *American Economic Journal: Microeconomics*, 10, 152–189. [1026]

Dai, Tianjiao and Juuso Toikka (2018), “Robust incentives for teams.” Unpublished paper, Department of Economics, MIT. [1027, 1042]

Deloitte (2016), “Global human capital trends: The new organization: Different by design.” <https://www2.deloitte.com/content/dam/Deloitte/global/Documents/HumanCapital/gx-dup-global-human-capital-trends-2016.pdf>. [1023]

Dong, Miaomiao (2018), “Strategic experimentation with asymmetric information.” Unpublished Paper, Pennsylvania State University. [1027]

Fleming, Wendell H. and Halil M. Soner (2006), *Controlled Markov Processes and Viscosity Solutions*, volume 25 of Stochastic Modelling and Applied Probability. Springer, Berlin. [1054]

Freixas, Xavier, Roger Guesnerie, and Jean Tirole (1985), “Planning under incomplete information and the ratchet effect.” *Review of Economic Studies*, 52, 173–191. [1025]

Fudenberg, Drew and Jean Tirole (1986), “A “signal-jamming” theory of predation.” *RAND Journal of Economics*, 17, 366–376. [1026]

Georgiadis, George (2015), “Projects and team dynamics.” *Review of Economic Studies*, 82, 187–218. [1027]

Georgiadis, George (2017), “Deadlines and infrequent monitoring in the dynamic provision of public goods.” *Journal of Public Economics*, 152, 1–12. [1027]

Gully, Stanley M., Kara A. Incalcaterra, Aparna Joshi, and J. Matthew Beaubien (2002), “A meta-analysis of team-efficacy, potency, and performance: Interdependence and level of analysis as moderators of observed relationships.” *Journal of applied psychology*, 87, 819–832. [1027]

Hackman, J. Richard (1987), “The design of work teams.” In *Handbook of Organizational Behavior* (Jay William Lorsch, ed.), 315–342, Prentice-Hall, Englewood Cliffs, New Jersey. [1023]

Halac, Marina, Navin Kartik, and Qingmin Liu (2017), “Contests for experimentation.” *Journal of Political Economy*, 125, 1523–1569. [1027]

Hermalin, Benjamin E. (1998), “Toward an economic theory of leadership: Leading by example.” *American Economic Review*, 88, 1188–1206. [1024, 1043]

Hölmstrom, Bengt (1982), “Moral hazard in teams.” *Bell Journal of Economics*, 13, 324–340. [1024, 1026]

Hölmstrom, Bengt (1999), “Managerial incentive problems: A dynamic perspective.” *Review of Economic Studies*, 66, 169–182. [1024, 1026, 1034]

Keller, Godfrey, Sven Rady, and Martin Cripps (2005), “Strategic experimentation with exponential bandits.” *Econometrica*, 73, 39–68. [1025, 1029]

Klein, Nicolas and Peter Wagner (2018), “Strategic investment and learning with private information.” Unpublished paper. [1027]

Lazear, Edward P. and Kathryn L. Shaw (2007), “Personnel economics: The economist’s view of human resources.” *Journal of economic perspectives*, 21, 91–114. [1023]

Legros, Patrick and Steven A. Matthews (1993), “Efficient and nearly-efficient partnerships.” *Review of Economic Studies*, 60, 599–611. [1026]

Liptser, Robert and Albert N. Shiryaev (2013), *Statistics of Random Processes: II. Applications*, volume 6 of Stochastic Modelling and Applied Probability. Springer, New York. [1052]

Marx, Leslie M. and Steven A. Matthews (2000), “Dynamic voluntary contribution to a public project.” *Review of Economic Studies*, 67, 327–358. [1026]

Mathieu, John, Travis M. Maynard, Tammy Rapp, and Lucy Gilson (2008), “Team effectiveness 1997–2007: A review of recent advancements and a glimpse into the future.” *Journal of Management*, 34, 410–476. [1027]

Olson, Mancur (1965), *The Logic of Collective Action: Public Goods and the Theory of Groups*. Harvard University Press, Cambridge, Massachusetts. [1026]

Prat, Julien and Boyan Jovanovic (2014), “Dynamic contracts when the agent’s quality is unknown.” *Theoretical Economics*, 9, 865–914. [1025]

Radner, Roy, Roger Myerson, and Eric Maskin (1986), “An example of a repeated partnership game with discounting and with uniformly inefficient equilibria.” *Review of Economic Studies*, 53, 59–69. [1026]

Riordan, Michael H. (1985), “Imperfect information and dynamic conjectural variations.” *RAND Journal of Economics*, 16, 41–50. [1026]

Teschl, Gerald (2012), *Ordinary Differential Equations and Dynamical Systems*, volume 140 of Graduate Studies in Mathematics. American Mathematical Society, Providence, Rhode Island. [1047]

Weitzman, Martin L. (1980), “The “ratchet principle” and performance incentives.” *Bell Journal of Economics*, 11, 302–308. [1025]

Whitt, Ward (1980), “Some useful functions for functional limit theorems.” *Mathematics of Operations Research*, 5, 67–85. [1054]

Yildirim, Huseyin (2006), “Getting the ball rolling: Voluntary contributions to a large-scale public project.” *Journal of Public Economic Theory*, 8, 503–528. [1026]

Co-editor Thomas Mariotti handled this manuscript.

Manuscript received 24 January, 2019; final version accepted 5 December, 2019; available online 17 December, 2019.