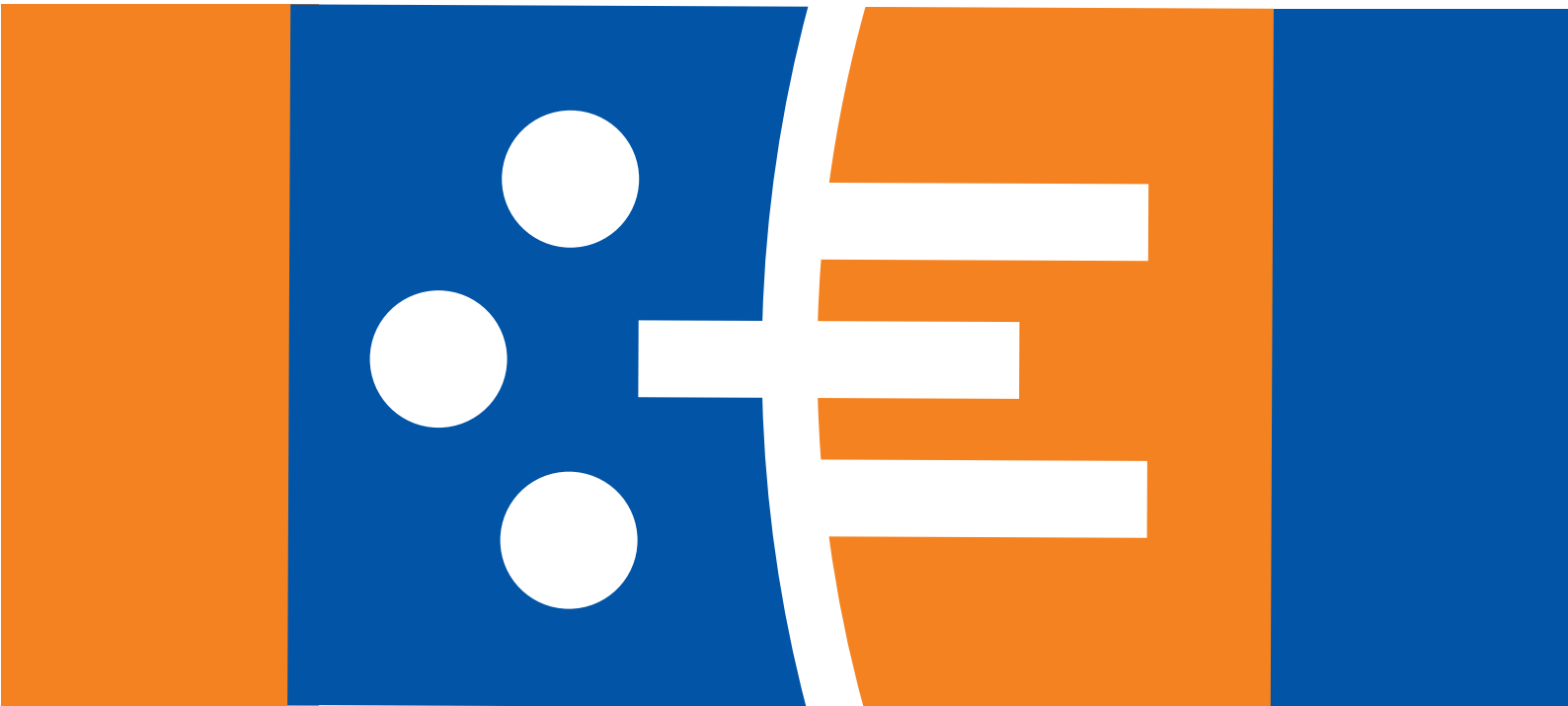




DPEB 24/01

Coordinated punishment? Only if reciprocators abound

- > **Adriana Alventosa**
Universitat de València and ERI-CES, Spain
- > **Vicente Calabuig**
Universitat de València and ERI-CES, Spain
- > **Gonzalo Olcina**
Universitat de València and ERI-CES, Spain

May, 2024

Coordinated punishment? Only if reciprocators abound

Adriana Alventosa[†], Vicente Calabuig^{†,*}, Gonzalo Olcina[†]

[†]ERI-CES, University of Valencia

*Corresponding author

May 7, 2024

Abstract

In this paper, we explore under which conditions coordinated punishment performs better than uncoordinated punishment, yielding higher levels of joint investment and boosting the return set by allocators in asymmetric situations as a team trust game. To this end, we theoretically compare these two different punishment schemes in a team trust game with information asymmetries: (i) an uncoordinated punishment system where individual punishment destroys the punished agent's payoff and (ii) a coordinated punishment system where it is necessary that the number of individuals willing to carry out punishment exceeds a particular threshold for the punishment to be effective. Our results reveal that only when there is a sufficiently high proportion of reciprocal investors in the population, a coordinated punishment system leads to efficient equilibria in a wider range of cases than uncoordinated punishment and it does so with higher returns. However, for intermediate proportions of reciprocators, uncoordinated punishment performs better.

Keywords: coordinated punishment, uncoordinated punishment, team trust game, efficiency

JEL classification: C72, D82, D90.

Authors' address: Av. Tarongers s/n, 46022, Valencia, Spain.

Acknowledgments: Authors acknowledge financial support from the Spanish Ministry of Economics, Innovation and Competitiveness under project ECO2017-87245 and grant BES-2015-073089, the Spanish Ministry of Science and Innovation under project PID2021-128228NB-100 and the Conselleria d'Innovació, Universitat, Ciència i Societat Digital (Generalitat Valenciana) under project AICO/2021/257.

1 Introduction

Many real-life economic situations can be represented as trust games with a team of investors. Team investment situations where the proceeds of aggregate investment of the team members (investors) are under the control of another agent (the allocator) are quite ubiquitous in real economies, especially, in financial and labour markets. Several investors (principals) delegate to a broker (agent) their decision about how to invest their funds and the broker may take advantage of them. Another prominent example appears in the labour market. In many employment relations, a group of employees is hired by a single employer (the firm). The labour contract in these cases is highly incomplete and usually assigns significant authority to the employer. This asymmetric distribution of decision rights puts the employees in danger of being exploited, leading to inefficiencies if they refuse to cooperate.

These inefficient outcomes have been approached in the literature in several ways, such as allowing for reputation building by the allocator, social preferences (a preference for fairness by the allocator) or, most commonly, endowing the investors (employees) with punishment power over the allocator (employer). However, punishment in groups entails an additional challenge: free riding. If we consider several investors deciding on punishing a single allocator, they could potentially have incentives to free ride and benefit from the costly punishment decision of another investor. Therefore, punishment becomes a public good itself with second-order free riding (Molleman, Kölle, Starmer and Gächter, 2019). This happens whenever punishment is *uncoordinated* in the sense that it is effective right from the first punisher.

A possible solution to the free-riding problem in punishment is to make punishment *coordinated*, that is, to require a minimum threshold in the number of punishers for punishment to be effective. Examples of the use of coordinated punishment when punishment has to be inflicted by the group abound in the historical evidence. Punishment can be implemented both formally, with penalty fees, or informally, through social punishment such as ostracism or hostility. Consider, for instance, the institution of ostracism in ancient Athens, the punishment of non-compliant-norm members by the village in rural areas in Africa or the expelling of free-riding members in medieval guilds. In modern societies, there are also many instances of coordinated punishment in smaller institutions such as partnerships, professional societies and unions. Furthermore, there is experimental evidence on coordinated punishment being preferred over uncoordinated punishment (Molleman et. al, 2019).

We are interested in understanding when, and under which conditions, will coordinated punishment perform better than uncoordinated punishment in the sense of yielding higher levels of joint investment and boosting the return set by allocators. We pursue characterising the conditions that make each system prevail and understanding whether this prevalence is correlated with the preference distribution of the involved agents.

According to some authors (Boyd et al. 2010, Olcina and Calabuig 2015), a possible reason for the better performance of coordinated punishment, is that it is associated with increasing returns to scale. This feature implies that the punishment effectiveness (i.e., the damage inflicted on the target) increases exponentially with the number of punishers. This, in turn, calls for modelling coordinated and uncoordinated punishment as two different technologies as it is implicitly assumed that individuals using coordinated punishment can inflict more damage at the same cost when they punish together as a group than when they do it individually. As a result, the benefits of coordinated punishment compared with uncoordinated punishment could be explained by a higher effectiveness of the punishment in the former device. However, in this work we consider constant returns to scale (coordinated punishment is as damaging as joint uncoordinated punishment), in order to analyse just the strategic component of coordinated punishment, that is, when and how it is more effective in solving the free-rider problem among investors. Our results are even stronger in the presence of increasing returns to scale, so we are in fact characterising the lower bound.

We theoretically compare two different punishment schemes with constant returns to scale in a team trust game with two investors and one allocator with information asymmetries: (i) an uncoordinated punishment system where individual punishment destroys the allocator's payoff and (ii) a coordinated punishment system where it is necessary that the number of investors willing to carry out punishment exceeds a particular threshold for the punishment to be effective.

Both types of players can either have standard preferences (selfish investors and selfish allocators) or social concerns (reciprocal investors and fair-minded allocators), being types private information. In the first stage of the game, the two investors individually and simultaneously decide whether or not to invest in a joint project. Only if both decide to invest, a positive surplus is generated and the team's allocator must decide how much of such surplus to keep for himself/herself and how much return to the investors. Once these investors observe their return, they can implement punishment. For the game with

uncoordinated punishment, the allocator has a proportion of his/her payoff destroyed as soon as one of the two investors decides to punish him/her. For the coordinated punishment case, however, it is necessary that both investors decide to punish the allocator for such a decision to have a negative impact on the allocator's payoff. In the labour context, players could resemble two employees (the investors) and one employer (the allocator), where employees decide whether to exert effort on a joint project (investment decision) or not and the project only succeeds if they both do so. If this is the case, the employer can decide how much of the project's profits pay to the employees (returns).¹ Employees, in turn, can punish by going on strike.

Our main result states that coordinated punishment performs better in terms of efficiency only if there are enough reciprocal investors in the population. That is, only when the proportion of reciprocators in the population is sufficiently high, a coordinated punishment system leads to equilibria with joint investment, positive rewards and no punishment in a wider range of cases than uncoordinated punishment. With intermediate proportions of reciprocators, however, uncoordinated punishment is superior to coordinated punishment, as positive rewards that are not punished are returned by selfish allocators under uncoordinated punishment while null rewards that are punished are returned under coordinated punishment. Finally, with low proportions of reciprocators, however, differences between coordinated and uncoordinated punishment are minimal.

The superiority of the coordinated scheme, which happens for high proportions of reciprocators, is driven by the fact that reciprocators are willing to punish higher offers under coordinated punishment, to which selfish allocators return higher rewards than under uncoordinated punishment in order to avoid punishment. This, in turn, enhances efficiency through joint investment even with a lower proportion of fair-minded allocators in the population.

This effect is indeed given by coordination as a mechanism to overcome the free-rider issue. Intuitively, differences between the two types of punishment lie in the free-riding incentives between reciprocators. With coordinated punishment, this behaviour can be avoided when there are many reciprocators in the population as if one of the two investors free rides, punishment actions will be costly and ineffective in destroying the allocator's payoff. When the proportion of reciprocators is high, the probability of being paired with

¹We assume that investment decisions are complements (Harrison and Hirschleifer, 1989; Van Huyck et al. 1990; Brandts and Cooper, 2006; Riedl et al. 2015; Calabuig and Olcina, 2015; Calabuig et al. 2022) and that the allocator values the investors' decisions equally and, therefore, returns the same amount to them in case of joint investment.

another reciprocator is also high, so the chance of there being joint punishment is higher. Uncoordinated punishment, however, is prone to free riding for any investors preference distribution, as the allocator’s payoff is going to be undermined even if there is only one punisher. Hence, even if the proportion of reciprocators is high, the threat of joint punishment is not as strong as with coordinated punishment.

1.1 Related Literature

To the best of our knowledge, the first paper dealing with coordinated punishment was by Boyd, Gintis and Bowles (2010), being pioneers in considering that punishment costs in a public goods game shall decrease with the number of punishers. Boyd et al. (2010) show how can punishment proliferate under these conditions. In this line, García and Traulsen (2019) find that indeed coordinated punishment can emerge from individual interactions, but the stability of the associated institutions over time is weak whenever the institutions implementing punishment cannot be observable by individuals. However, these works represent symmetric situations, public goods games, while our aim is to understand prevailing asymmetric situations of real-life economies like principal-agent relationships.

This work is also closely related to that by Calabuig and Olcina (2015), a team trust game with coordinated punishment.² Their main result is that cooperation only evolves and is maintained if there is enough punishment capacity in society and there are enough individuals willing to implement punishment. The fully cooperative equilibrium is achieved under the threat of effective coordinated punishment, which is supported by the presence of a high proportion of reciprocators in the population, in line with our results. However, the authors only provide results for coordinated punishment. Our goal in this work is to compare coordinated punishment with uncoordinated punishment in terms of joint investment, returns and punishment actions, in order to highlight under which conditions is one system better than the other. Furthermore, we consider constant returns to scale in punishment rather than increasing returns to scale in order to isolate the effect of the strategic component of coordinated punishment.

Experimental evidence on the comparison between coordinated and uncoordinated punishment in team trust games exists (Calabuig, Jiménez, Olcina and Rodríguez-Lara, 2022), showing that joint investment and returned rewards are enhanced when punish-

²In their work, they also introduce a peer punishment stage at the end of the game, which we do not consider here.

ment is coordinated. The equilibria are characterised for a particular set of values. Our work, in contrast, entails a generalisation of this comparison revealing that such superiority is only given for a set of values and that, otherwise, uncoordinated punishment can be better at sustaining efficient equilibria.

Other experimental works have highlighted the preference of individuals to either tacitly coordinate on punishment even with second-order free riding incentives (Diekmann and Przepiorka, 2015), or to choose a coordinated punishment scheme when given the chance to choose between coordinated and uncoordinated punishment (Molleman et al., 2019). This suggests that coordinated punishment may not only be strategically superior but also preferred by individuals. Our contribution is to understand under which conditions it is indeed more efficient.

The rest of this paper is organized as followed. In section 2 we present the model to analyse, describing as well the different preferences considered. Section 3 collects the results of the model, stage by stage and, in the last place, Section 4 presents concluding remarks.

2 The model

2.1 Team trust game with punishment

Consider a 3-player team In this paper, we explore under which conditions coordinated punishment performs better than uncoordinated punishment, yielding higher levels of joint investment and boosting the return set by allocators in asymmetric situations as a team trust game. To this end, we theoretically compare these two different punishment schemes in a team trust game with information asymmetries: (i) an uncoordinated punishment system where individual punishment destroys the punished agent’s payoff and (ii) a coordinated punishment system where it is necessary that the number of individuals willing to carry out punishment exceeds a particular threshold for the punishment to be effective. Our results reveal that only when there is a sufficiently high proportion of reciprocal investors in the population, a coordinated punishment system leads to efficient equilibria in a wider range of cases than uncoordinated punishment and it does so with higher returns. However, for intermediate proportions of reciprocators, uncoordinated punishment performs better. trust game with two investors and one allocator.

In the first stage of this game, *the investment stage*, the two investors simultaneously and independently decide whether to invest (I) or not invest (NI) in a project which can generate a surplus of 2 only if both of them decide to carry out the investment. In other words, only the combination of actions (I, I) leads to the project being successful but if at least one investor decides not to invest, the game ends and all players obtain a payoff of 0.³ Investing is a costly action for investors. In particular, it entails a cost $k \in (0, \frac{1}{2})$ regardless of the success or failure of the project. Not investing, on the other side, is costless.

In the second stage, *the reward stage*, the allocator decides how much of the surplus generated he/she is willing to return to the investors. With this end, the allocator decides the payoff b , the reward to each of the investors ($0 \leq b \leq 1$), being this symmetric to both of them.⁴ Recall this stage is only reached if both investors decide to invest in the first stage.

In the third stage, *the punishment stage*, the investors observe their reward b for their investment and simultaneously and independently decide whether to punish (p) or not to punish (np) the allocator. We consider that punishing has a cost of z/n for each punishing investor, where z is the total cost and n is the number of investors who punish. Not punishing, on the other side, is costless.

The impact of punishment on the allocator's payoff depends on the punishment scheme we consider. Under a coordinated punishment scheme, only if both investors coordinate to punish, the allocator will get his/her payoff destroyed in a proportion λ_c , where $0 \leq \lambda_c \leq 1$. Under an uncoordinated punishment scheme, however, the individual punishment action has an impact of λ_u on the allocator's payoff, where $0 \leq \lambda_u \leq 1$, regardless of whether the investors coordinate in punishing or not. If they do so, the impact will be $2\lambda_u$ instead. Table 1 summarises these details.

	Coordinated	Uncoordinated
p, p	λ_c	$2 \lambda_u$
p, np or np, p	0	λ_u
np, np	0	0

Table 1: Impact of punishment

³We focus on a symmetric situation where investments are considered complements in order to cleanly compare coordinated and uncoordinated punishment. Considering scenarios where only one investor invests and the project is still carried out would not give any added value to the analysis.

⁴In this paper, we focus on symmetric equilibria.

Given that only if both investors invest, the game goes on to the reward and punishment stage, in the following we will only present the final payoff functions for the (I, I) subgame.

The material payoffs for each of the investors will be determined by:

$$\pi_I = b - k - z/n \quad \text{with punishment} \quad (1)$$

The material payoffs of the allocator will depend on the punishment scheme, which can either be coordinated (CP) or uncoordinated (UP):

$$\pi_A(CP) = \begin{cases} 2(1-b)(1-\lambda_c) & \text{if both punish} \\ 2(1-b) & \text{otherwise} \end{cases} \quad (2)$$

$$\pi_A(UP) = \begin{cases} 2(1-b)(1-2\lambda_u) & \text{if both punish} \\ 2(1-b)(1-\lambda_u) & \text{only one punishes} \\ 2(1-b) & \text{if nobody punishes} \end{cases} \quad (3)$$

Suppose there is complete information and that all players have selfish preferences. By backward induction, selfish investors will never implement costly punishment, np , the selfish allocator will keep all of the surplus for himself/herself and return $b = 0$ and investors, anticipating the allocator's behaviour, will decide not to carry out the costly investment (NI, NI) . This, however, will lead to an inefficient Perfect Equilibrium where every player receives a payoff of 0.

2.2 Social preferences

Experimental evidence has consistently proved that individual preferences only follow the standard selfishness assumption in a limited number of cases (Fehr and Schmidt, 1999; Fehr and Gächter, 2000). For this reason, we consider it necessary to introduce social preferences into our model.

On the side of the investors, we consider a certain proportion q of the population from which the two investors are drawn to be reciprocal investors, where we understand as reciprocal the willingness to punish hostile behaviour. That is, if the allocator returns an unfair offer, a reciprocal investor will implement punishment in the last stage. Given that

the surplus that each investor generates is of size 1, a reciprocal investor will consider fair a return of $1/2$ and unfair any return below $1/2$.⁵ We capture the disutility derived from unfair returns with a parameter $\alpha = 2$, measuring the proportional distance from the allocator's payoff to the investor's reward. This approach follows Fehr and Schmidt (1999) for disadvantageous inequality.⁶ The remaining $1 - q$ investors of the population will be considered selfish, with the payoff function formerly proposed.

The payoff function for a reciprocal investor in the subgame (I, I) if he/she receives a reward $b < 1/2$, will be determined by:

$$\pi_I(CP) = \begin{cases} b - k - z/2 - \alpha \max\{0, (1 - b)(1 - \lambda_c) - b\} & \text{if both punish} \\ b - k - z - \alpha \max\{0, (1 - b) - b\} & \text{if only he/she punishes} \\ b - k - \alpha \max\{0, (1 - b) - b\} & \text{otherwise} \end{cases} \quad (4)$$

$$\pi_I(UP) = \begin{cases} b - k - z/2 - \alpha \max\{0, (1 - b)(1 - 2\lambda_u) - b\} & \text{if both punish} \\ b - k - z - \alpha \max\{0, (1 - b)(1 - \lambda_u) - b\} & \text{if only he/she punishes} \\ b - k - \alpha \max\{0, (1 - b)(1 - \lambda_u) - b\} & \text{if only the other one punishes} \\ b - k - \alpha \max\{0, (1 - b) - b\} & \text{if nobody punishes} \end{cases} \quad (5)$$

On the side of the allocator, we consider a certain proportion m of the population from which the allocator is drawn to be fair-minded allocators, where we understand as fair-minded having as a dominant strategy returning the fair reward $b = 1/2$. The remaining $(1 - m)$ are selfish.

2.3 Assumptions

Before moving on to resolving the sequential game by backward induction, we need to make three assumptions on the relationship between the parameters that characterize the punishing institutions.

⁵The analysis could be analogously done with any other reference point different from $b = 1/2$.

⁶We assume that investors do not have a disutility for generous rewards above $1/2$, which reflects Fehr and Schmidt's (1999) advantageous inequality.

Assumption 1: $2\lambda_u = \lambda_c$

Assumption 1 states that if both investors punish, uncoordinated punishment is as destructive for the allocator as coordinated punishment. As it is common in the literature, we could alternatively assume that coordinated actions can lead to synergies and economies of scale such that if both investors punish, coordinated punishment is *more* destructive than uncoordinated punishment, i.e. $2\lambda_u < \lambda_c$ (Boyd et al., 2010; Calabuig and Olcina, 2015). Considering this kind of returns would strengthen our results even further. However, we want to analyse the advantages of coordinated punishment vs. uncoordinated punishment with the same returns. Therefore, we will assume that $2\lambda_u = \lambda_c$.

Assumption 2: $2\lambda_u(\frac{1}{2-2\lambda_u}) \leq z \leq 2\lambda_u$

The upper bound is necessary to guarantee that the behaviour of a reciprocal investor differs from the behaviour of a selfish one. Briefly, in uncoordinated punishment, a reciprocal investor that punishes reduces his/her disadvantageous inequality with respect to the allocator in $\alpha\lambda_u = 2\lambda_u$. Thus, this effect must compensate the cost z of punishing the allocator, such that he/she indeed does so. Notice that, by assumption 1, if $z \leq 2\lambda_u$, then $z \leq 2\lambda_c$ holds. The lower bound, on the other hand, is a simplifying assumption in order to avoid numerous non-interesting subcases arising. Besides that, it is plausible to impose a lower bound to z , as punishment must be a costly action. As the destructive power of punishment increases, so does the lower bound. This reflects how highly destructive institutions may face high punishment costs.

Assumption 3: $\lambda_c \geq 1/2$

Finally, assumption 3 is necessary for coordinated punishment to have a sufficient impact on the behaviour of the allocator. If $\lambda_c < 1/2$, the allocator would prefer to offer a reward of $b = 0$ instead of a fair reward of $b = 1/2$, even if he/she knew that the investors were going to punish him/her. Notice that assumptions 1 and 3 imply $1/4 \leq \lambda_u \leq 1/2$.

3 Results

In this section, we derive the Perfect Bayesian Equilibria (PBE) of the two sequential games: the team trust game with coordinated punishment and the team trust game with uncoordinated punishment. To do so, we will characterize behaviour in each subgame,

anticipating what will happen in the following stages. In order to compare the potential superiority of one type of punishment scheme over the other one, we will focus on the efficient pooling equilibria where both types of investors decide to invest, both types of allocators return positive rewards and there is no punishment by the investors.

3.1 Punishment stage

In this last stage, investors must decide whether to punish the allocator for the reward b received. Given that types are private information, we define μ as the updated probability of facing a reciprocal investor after observing (I, I) in the investment stage. That is, $\mu = \text{Prob}(r/(I, I))$, where r stands for a reciprocal type. Any punishment subgame is characterized by a belief μ and a reward b , thus, we denote this subgame by $CP(\mu, b)$ for the coordinated punishment scenario and $UP(\mu, b)$ for the uncoordinated one. We represent the symmetric Bayesian Nash Equilibria (BNE) of this subgame by profiles (x, y) , where the first term represents the action of the reciprocal type and the second the action of the selfish type.

First, notice that if the allocator returns a fair or generous reward $b \geq 1/2$, no type of investor will punish. In the following lemmata, we present the solution of $CP(\mu, b)$ and $UP(\mu, b)$ for unfair rewards ($b < 1/2$). Formal proofs have been relegated to Appendix A.

Lemma 1: *The solution of any subgame $CP(\mu, b)$ with unfair rewards is:*

- a) If $0 \leq b < b_1^c$, then (p, np) if $\mu > \bar{\mu}$ and (np, np) otherwise
- b) If $b_1^c \leq b < b''$, then (p, np) if $\mu > \bar{\bar{\mu}}$ and (np, np) otherwise
- c) If $b'' \leq b < 1/2$, then $(np, np) \forall \mu$

$$\text{where } b_1^c = \frac{1-\lambda_c}{2-\lambda_c}, b'' = \frac{1}{2} - \frac{z}{8}, \bar{\mu} = \frac{z}{z/2+2\lambda_c(1-b)} \text{ and } \bar{\bar{\mu}} = \frac{z}{z/2+2(1-2b)}$$

Lemma 2: *The solution of any subgame $UP(\mu, b)$ with unfair rewards is:*

- a) If $0 \leq b < \hat{b}$, then $(p, np) \forall \mu$
- b) If $\hat{b} \leq b < b_1^u$, then (p, np) if $\mu > \tilde{\mu}$ and (np, np) otherwise
- c) If $b_1^u \leq b < b^{**}$, then (p, np) if $\mu > \tilde{\tilde{\mu}}$ and (np, np) otherwise

d) If $b^{**} \leq b < 1/2$, then $(np, np) \forall \mu$

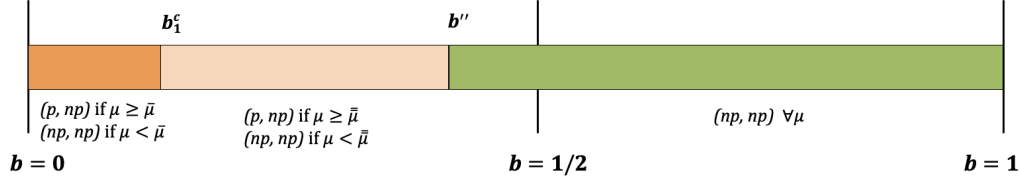
$$\text{where } \hat{b} = 1 - \frac{z}{2\lambda_u}, b_1^u = \frac{1-2\lambda_u}{2-2\lambda_u}, b^{**} = \frac{2(1-\lambda_u)-z/2}{2(2-\lambda_u)}, \tilde{\mu} = \frac{2z-4\lambda_u(1-b)}{z} \text{ and } \tilde{\tilde{\mu}} = \frac{z-2\lambda_u(1-b)}{z/2-4\lambda_u(1-b)+2(1-2b)}$$

Selfish investors have as a dominant strategy not to punish, whereas reciprocal investors may have incentives to do so. Notice that for an intermediate range of values of the reward, reciprocal investors need that beliefs of being paired with another reciprocal investor are high enough in order to punish. Otherwise, it does not compensate to undertake the punishment costs by themselves. Additionally, as rewards increase, even though they do not reach the fair reward of $b = 1/2$, reciprocators stop punishing if they receive a high enough reward for their investment. These results hold for both types of punishment systems.

Under uncoordinated punishment, however, there is also a lower reward segment ($0 \leq b < \hat{b}$) which reciprocal investors punish independently of their beliefs. In other words, the reward is so low that they punish it because by doing so they can reduce inequality even if they are the only punisher. This, nevertheless, does not happen with coordinated punishment given that individual punishment does not lead to a destruction in the allocator's payoffs. Thresholds on the beliefs are increasing in b and z and decreasing in λ_i . Formal proof has been relegated to the appendix.

If we compare these two situations we can shed light on the frequency of punishment in each context (see Figure 1). We can synthesise lemma 1 by saying that, under coordinated punishment, rewards $0 \leq b < b''$ are punished if there are enough reciprocators, whereas rewards $b'' \leq b \leq 1$ are left unpunished. In the same way, under uncoordinated punishment (lemma 2), if $0 \leq b < \hat{b}$, there is unconditional punishment, rewards $\hat{b} \leq b < b^{**}$ are conditionally punished and rewards $b^{**} \leq b < 1$ are always left unpunished.

Coordinated Punishment



Uncoordinated Punishment

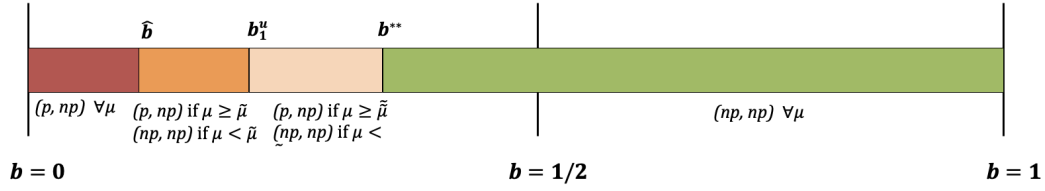


Figure 1: BNE of the punishment stage

While the relation among b_1^c , $\hat{b}$ and b_1^u is ambiguous, $b'' > b^{**}$ always holds. This implies that the range of values of b that reciprocal investors leave unpunished is larger under uncoordinated punishment. In other words, coordinated punishers exert a stronger punishment pressure by expecting higher rewards from the allocators.

Corollary 1: *For a sufficiently high proportion of reciprocators in the population, punishment occurs more frequently under a coordinated punishment system than under an uncoordinated one.*

3.2 Reward stage

This reward stage is only reached in the subgame (I, I) of the investment stage. Thus, after observing that both investors have decided to invest in the project, the allocator chooses how much to return to each investor, $0 \leq b \leq 1$, which we assume symmetric. As the project surplus is of size 2, each investor generates a surplus of 1 with his/her investment, so we define a reward of $b = 1/2$ as fair. Recall that we also consider two possible different types of allocators, fair-minded⁷ who will return the fair reward $b_f^i = 1/2$ and selfish who will return b_s^i . The super index $i \in \{c, u\}$ represents the punishment scheme: coordinated (c) or uncoordinated (u). It should be noted that fair rewards are never going

⁷Notice that a fair-minded allocator does not change his/her rewarding policy when there is incomplete information.

to be punished by any type of investor. Unfair rewards, that is, any $b < 1/2$, may or may not be punished by reciprocators.

When deciding on how much to reward, a selfish allocator will compare the expected cost of being punished with the expected cost of avoiding it. On the one hand, the expected cost of punishment will depend on the beliefs the allocator has on the types of the investors. On the other hand, the expected cost of avoiding punishment is the minimum reward an allocator shall return for no investor to punish.

We describe the optimal reward policy of the selfish allocator with the following set of lemmata where the thresholds $\bar{\mu}(b)$ and $\tilde{\mu}(b)$ follow from Lemmas 1 and 2 evaluated at different rewards b . Formal proof has been relegated to Appendix A.

Lemma 3: *In the team trust game with coordinated punishment, the reward policy of a selfish allocator is:*

- a) If $\mu < \bar{\mu}(0)$ then $b_s^c = 0$
- b) If $\bar{\mu}(0) \leq \mu < \bar{\mu}(b_1^c)$, then $b_s^c = 0$ if $f(\mu) > 0$ or $b_s^c = b^c$ if $f(\mu) \leq 0$
- c) If $\bar{\mu}(b_1^c) \leq \mu \leq 1$, then $b_s^c = 0$ if $g(\mu) > 0$ or $b_s^c = b^{cc}$ if $g(\mu) \leq 0$

where: $b^c = \frac{\mu(z/2+2\lambda_c)-z}{2\mu\lambda_c}$ and $b^{cc} = \frac{\mu(z/2+2)-z}{4\mu}$, $f(\mu) = -(2\lambda_c^2)\mu^3 + (2\lambda_c + z/2)\mu - z$ and $g(\mu) = -(4\lambda_c)\mu^3 + (2 + z/2)\mu - z$

Lemma 4: *In the team trust game with uncoordinated punishment, the reward policy of a selfish allocator is:*

- a) If $\mu < \tilde{\mu}(\hat{b})$ then $b_s^u = 0$
- b) If $\tilde{\mu}(\hat{b}) \leq \mu < \tilde{\mu}(b_1^u)$ then $b_s^u = b^u$
- c) If $\tilde{\mu}(b_1^u) \leq \mu \leq 1$, then $b_s^u = 0$ if $h(\mu) > 0$ or $b_s^u = b^{uu}$ if $h(\mu) \leq 0$.

where: $b^u = \frac{2\lambda_u + \mu z/2 - z}{2\lambda_u}$, $b^{uu} = \frac{\mu(z/2+2(1-2\lambda_u))-z+2\lambda_u}{2(\lambda_u(1+2\mu)+2\mu)}$ and $h(\mu) = -\mu^2(8\lambda_u(1 + \lambda_u)) + \mu(z/2 + 2(1 - 2\lambda_u - 2\lambda_u^2)) - z + 2\lambda_u$

If the proportion of reciprocal investors is sufficiently low, a selfish allocator anticipates that there is going to be no punishment in the last stage, so the expected costs of

being punished are zero. Thus, the selfish allocator prefers to offer the lowest return $b_s^i = 0$.

However, as the proportion of reciprocal investors increases, so does the probability of being punished. This makes selfish allocators consider the possibility of either returning nothing and facing punishment or, alternatively, offering a positive minimal reward (b^c, b^{cc}, b^u or b^{uu} as the case may be) such that reciprocal investors do not punish them. The expected costs of being punished are $\lambda_c \mu^2$ under coordinated punishment and $2\lambda_u \mu$ under uncoordinated punishment, where the difference is that coordinated punishment requires two reciprocators for punishment to be effective. Therefore, the optimal reward policy depends on the critical values derived from the comparison of the expected payoffs by offering $b_s^i = 0$ and $b_s^i > 0$.

The comparison between the different rewarding policies leads to a cubic equation $f(\mu) = -(2\lambda_c^2)\mu^3 + (2\lambda_c + z/2)\mu - z$ for the comparison between offering $b_s^c = 0$ or $b_s^c = b^c$ when $\bar{\mu}(0) \leq \mu < \bar{\mu}(b_1^c)$; a cubic equation $g(\mu) = -(4\lambda_c)\mu^3 + (2 + z/2)\mu - z$ for the comparison between offering $b_s^c = 0$ or $b_s^c = b^{cc}$ when $\bar{\mu}(b_1^c) \leq \mu \leq 1$; and a quadratic equation $h(\mu) = -(8\lambda_u(1 - \lambda_u))\mu^2 + (z/2 + 2(1 - 2\lambda_u(1 - \lambda_u)))\mu - z + 2\lambda_u$ for the comparison between offering $b_s^u = 0$ or $b_s^u = b^{uu}$ when $\tilde{\mu}(b_1^u) \leq \mu \leq 1$.

Intuitively, when the proportion of reciprocal investors in the population surpasses a certain threshold ($\bar{\mu}(0)$ in coordinated punishment and $\tilde{\mu}(b^u)$ in uncoordinated punishment), reciprocal investors may start considering punishing if the return they receive is too low in order to reduce payoff inequality. The higher the proportion of these investors, the higher the return necessary to avoid punishment ($\frac{\partial b^c}{\partial \mu}, \frac{\partial b^{cc}}{\partial \mu}, \frac{\partial b^u}{\partial \mu}$ and $\frac{\partial b^{uu}}{\partial \mu} > 0$). For certain values of μ , the required b_s^i is not too high and allocators prefer to reward investors and avoid punishment because it compensates them. However, as μ increases, so does b_s^i and the cost of avoiding punishment may become too high with respect to the cost of facing it. In these cases, they return nothing. The characterization of these values is found in Appendix A.

3.3 Investment stage: Perfect Bayesian Equilibria

Finally, in the investment stage, investors foresee the allocator's reward policy and their punishment reaction when deciding whether to invest or not. There are multiple Perfect Bayesian Equilibria (PBE) of the game with coordinated and uncoordinated punishment. For instance, there is always an inefficient non-cooperative equilibrium where both types

of investors choose not to invest for every q and m , as well as separating equilibria.⁸

However, in this work, we are going to focus on pooling equilibria where both types of investors decide to invest and, in particular, in the pooling equilibria where the selfish allocator decides to return a positive (but unfair) reward, the fair-minded allocator returns the fair reward and no type of investor punishes. We will call these equilibria *Positive Return Pooling Equilibria (PRPE)*. In these equilibria, $q = \mu = \text{Prob}(r/I, I)$ and there is no punishment. For completeness, other pooling PBE where both investors invest but selfish allocators return null rewards can be found in Appendix B.

The following propositions characterise the distributions of q and m of the PRPE under coordinated and uncoordinated punishment, and the reward policy of selfish allocators in these equilibria. Formal proof has been relegated to Appendix A.

Proposition 1: *In the team trust game with coordinated punishment, there exists a positive return pooling PBE if:*

- a) If $\bar{q}(0) \leq q < q_1$ or $q_2 \leq q < \bar{q}(b_1^c)$, and $m \geq m^c$, then $b_s^c = b^c$.
- b) If $\bar{q}(b_1^c) \leq q \leq q_1'$ or $q_2' \leq q < 1$, and $m \geq m^{cc}$, then $b_s^c = b^{cc}$.

where: $m^c = \frac{k-b^c+2(1-2b^c)}{0.5-b^c+2(1-2b^c)}$, $m^{cc} = \frac{k-b^{cc}+2(1-2b^{cc})}{0.5-b^{cc}+2(1-2b^{cc})}$, q_1 and q_2 are the positive roots of $f(\mu)$, and q_1' and q_2' are the positive root of $g(\mu)$.

Proposition 2: *In the team trust game with uncoordinated punishment, there exists a positive return pooling PBE if:*

- a) If $\tilde{q}(b^u) \leq q < \tilde{q}(b_1^u)$ and $m \geq m^u$, then $b_s^u = b^u$.
- b) If $q_1'' \leq q < 1$ and $m \geq m^{uu}$, then $b_s^u = b^{uu}$.

where $m^u = \frac{k-b^u+2(1-2b^u)}{0.5-b^u+2(1-2b^u)}$, $m^{uu} = \frac{k-b^{uu}+2(1-2b^{uu})}{0.5-b^{uu}+2(1-2b^{uu})}$ and q_1'' is the positive root of $h(\mu)$.

It can be proved that $b^{cc} > b^c$ and that $b^{uu} > b^u$, as well as $m^{cc} < m^c$ and $m^{uu} < m^u$. Therefore, the PRPE can be classified into two types. On the one side, there is a set of equilibria where a low reward is returned, b^c and b^u (Proposition 1a and 2a), and for which a high proportion of fair-minded allocators is needed, m^c and m^u . On the other

⁸This game has separating equilibria in which selfish investors invest and reciprocal investors do not invest, but these are not stable in long-run dynamics, following Calabuig and Olcina (2015).

side, for higher levels of q , both coordination schemes also have PRPE with high rewards, b^{cc} and b^{uu} , which will be returned even with a lower proportion of fair-minded allocators in the population, m^{cc} and m^{uu} (Proposition 1b and 2b).

Notice that for these PRPE to happen: (i) the proportion of reciprocal investors must be high enough and (ii) the proportion of fair-minded allocators must also be high enough. The first condition ensures that a selfish allocator sets a positive reward in order to avoid the potential punishment of reciprocal investors. The second condition, on the other side, guarantees that both types of investors prefer to invest rather than not to do so.

If these two conditions are met, both types of investors invest, regardless of their type. They receive a fair reward in case of being matched with a fair-minded allocator (who has as a dominant action to set these rewards) or a positive reward in case of being matched with a selfish allocator, following lemmas 3 and 4. If the investor is selfish, he/she will never punish regardless of the return received, whereas a reciprocal investor would not punish given that he/she either receives a fair reward or a sufficiently high one, by lemmas 1 and 2.

In order to obtain the critical proportion of fair-minded allocators needed to sustain this type of equilibrium, we obtain, for both types of investors, under which conditions will they invest, anticipating the positive reward that will be returned. Formal proof can be found in Appendix A. We find that as the proportion of reciprocal investors in the population increases, the proportion of fair-minded allocators required to sustain an efficient pooling PBE falls ($\frac{\partial m^j}{\partial q} < 0$ for $j \in \{c, cc, u, uu\}$) and the reward needed to avoid punishment increases ($\frac{\partial b^j}{\partial q} > 0$). This implies that as the population of reciprocators increases, investments with higher rewards that avoid punishment become more frequent.

3.4 Main Results

Our objective is, as previously stated, to compare both punishment systems in terms of achieving greater levels of efficiency. We present our main result concerning this issue in the following proposition. Formal proof has been relegated to Appendix A.

Proposition 3: *In the team trust game with punishment, for a proportion of reciprocal investors $q > q'_2$, the reward returned by a selfish allocator is higher under coordinated punishment, $b^{cc} \geq b^{uu}$ and the proportion of fair-minded allocators necessary for this to*

be a positive return pooling equilibrium is lower, $m^{cc} \leq m^{uu}$.

This proposition illustrates that in a population of investors with a high proportion of reciprocators, coordinated punishment is substantially superior to uncoordinated punishment. On the one hand, investors obtain higher rewards from allocators and, on the other hand, joint investment is achieved for a wider range of the allocators' preferences distribution.

In contrast with Proposition 3, the following proposition compares coordinated and uncoordinated punishment when the proportion of reciprocators is not high enough to achieve the PRPE under both schemes.

Proposition 4: *In the team trust game with punishment, for a proportion of reciprocal investors $q_1'' < q \leq q_2'$, the reward returned by a selfish allocator is higher under uncoordinated punishment, $b_s^c = 0$ and $b^{uu} > 0$ and the proportion of fair-minded allocators necessary for this to be a positive return pooling equilibrium is lower, $m^{uu} \leq m^{cp}$.*

For intermediate proportions of reciprocal investors, coordinated punishment would no longer be superior to uncoordinated punishment. Under coordinated punishment, zero rewards would be returned by selfish allocators and these would be punished. On the other hand, under uncoordinated punishment, positive rewards are always returned and investors never punish. Furthermore, the proportion of fair-minded allocators required for joint investment to happen is higher under coordinated punishment. Notice that for this range of q there is a PRPE under uncoordinated punishment but a null return pooling equilibrium with punishment under coordinated punishment. This set of equilibria, along with the corresponding threshold $m^{cp} = \frac{2+k+z-q(z/2+2\lambda_c)}{2+0.5+z-q(z/2+2\lambda_c)}$, can be found in Appendix B.

Finally, for sufficiently low proportions of reciprocal investors, there is no PRPE under any punishment scheme, and when there is, it happens for a very narrow range of q and with very low returns. For this reason, we claim that coordinated and uncoordinated punishment do not show substantial differences.

Remark: *In the team trust game with punishment, for a proportion of reciprocal investors $q \leq q_1''$, there are no substantial differences between coordinated and uncoordinated punishment.*

Numerical example

Let us illustrate the results using a numerical example. Suppose the parameters of the model take the following feasible values: $k = 0.25$, $z = 0.54$, $\lambda_c = 0.7$ and $\lambda_u = 0.35$, which fulfill the model's assumptions. Table 2 shows, for the different proportions of reciprocal investors, q , the necessary minimal fair-minded allocators proportion, m , for the existence of a pooling PBE where both investors invest, the corresponding reward by the selfish allocators, b_s and the punishment strategy of a reciprocator knowing the selfish investor never punishes, (x, np) . We also add a column indicating whether the PBE is the PRPE characterised in Propositions 1 and 2 or not. Recall that, despite we focus on equilibria where both investors invest, positive rewards are returned and there is no punishment (PRPE), there are other pooling PBE which, for completeness, we present in Appendix B.

This table points out that there is an inverse relationship between q and m . In other words, the higher the proportion of reciprocators in the population is, the lower the proportion of fair-minded allocators needed for there to be a pooling equilibrium in investment, regardless of the punishment scheme.

It is only for high values of q where the superiority of coordinated punishment can be appreciated. For this particular example, when the proportion of reciprocal investors is greater than or equal to 74.2%. When this happens, rewards under coordinated punishment (b^{cc}) are higher than under uncoordinated punishment (b^{uu}) and the proportion of fair-minded allocators is significantly lower under coordinated punishment (m^{cc}) than under uncoordinated punishment (m^{uu}). In fact, m drastically falls in this range, but only under coordinated punishment. Under uncoordinated punishment, a very high proportion of fair-minded allocators are still needed to sustain this type of equilibrium.

Nevertheless, for intermediate values of q , uncoordinated punishment dominates coordinated punishment, as for $q > 0.4187$ selfish allocators start returning positive rewards that stay unpunished under uncoordinated punishment. Under coordinated punishment, however, this does not happen until $q > 0.742$.

For low values of q , there is almost no difference between coordinated and uncoordinated punishment. Even if there is a range of values of q where an efficient pooling PBE appears under coordinated punishment, it happens in a very small segment ($0.33 \leq q \leq 0.348$) and with very low rewards. These rewards are b^c . Notice that b^u , on the side of uncoordinated punishment, seems to have disappeared. It is not the case, but

q	m		b_s		(x,np)		PRPE	
	CP	UP	CP	UP	CP	UP	CP	UP
0.05	0.9	0.8892	0	0	np	p	No	No
0.1	0.9	0.8850	0	0	np	p	No	No
0.15	0.9	0.8804	0	0	np	p	No	No
0.20	0.9	0.8754	0	0	np	p	No	No
0.25	0.9	0.8700	0	0	np	p	No	No
0.3	0.9	0.8641	0	0	np	p	No	No
0.3233	0.9	0.8611	0	0	np	p	No	No
0.33	0.8950	0.8602	0.024	0	np	p	Yes	No
0.34	0.8868	0.8589	0.058	0	np	p	Yes	No
0.348	0.8797	0.8579	0.084	0	np	p	Yes	No
0.35	0.8982	0.8576	0	0	p	p	No	No
0.4009	0.8945	0.8503	0	0	p	p	No	No
0.4187	0.8932	0.8025	0	0.2469	p	np	No	Yes
0.43	0.8923	0.8023	0	0.2470	p	np	No	Yes
0.45	0.8908	0.8021	0	0.2473	p	np	No	Yes
0.5	0.8866	0.8017	0	0.2479	p	np	No	Yes
0.55	0.8822	0.8012	0	0.2484	p	np	No	Yes
0.6	0.8773	0.8009	0	0.2489	p	np	No	Yes
0.65	0.8771	0.8005	0	0.2493	p	np	No	Yes
0.7	0.8664	0.8003	0	0.2497	p	np	No	Yes
0.742	0.5631	0.8000	0.386	0.2500	np	np	Yes	Yes
0.75	0.5556	0.8000	0.388	0.2500	np	np	Yes	Yes
0.8	0.5062	0.7998	0.399	0.2503	np	np	Yes	Yes
0.85	0.4525	0.7996	0.409	0.2506	np	np	Yes	Yes
0.9	0.3939	0.7994	0.418	0.2508	np	np	Yes	Yes
0.95	0.3298	0.7992	0.425	0.2510	np	np	Yes	Yes
1	0.2593	0.7990	0.433	0.2512	np	np	Yes	Yes

Table 2: PBE of the team trust game with (I,I), $k = 0.25$, $z = 0.54$, $\lambda_u = 0.35$, $q_1'' = 0.4187$ and $q_2' = 0.7420$.

it cannot be appreciated because it happens for very small values, i.e. $q < 0.0057$.

In the following Figures, we represent the PBE rewards b_s and minimal proportion of fair-minded allocators, m , for $q > q_1''$ using the values proposed in the numerical example, such that $q_1'' = 0.4187$. Recall that for $q \leq q_1''$, no substantial differences are observed between coordinated and uncoordinated punishment, so we graphically show when is coordinated punishment superior to uncoordinated punishment.

As it can be appreciated in Figure 2, rewards under uncoordinated punishment are

almost steady around $b^{uu} = 0.25$ over q . However, under coordinated punishment, rewards display a significant jump at $q = q'_2 = 0.742$. Namely, for $q < q'_2$ rewards are $b^c = 0$ (not PRPE), but at $q = q'_2$ rewards jump to an increasing pattern starting with $b^{cc} \sim 0.4$. This shows that for a sufficiently high proportion of reciprocal investors, coordinated punishment performs better in terms of rewards.

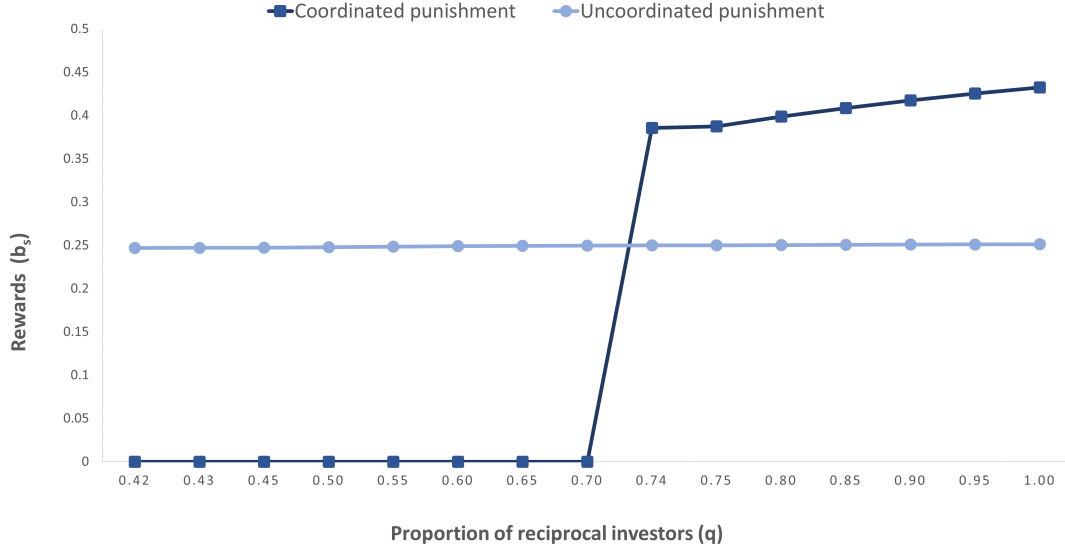


Figure 2: PBE rewards of the numerical example for $q > q''_1$.

Analogously, in Figure 3 we show that the proportion of fair-minded allocators needed for there to be joint investment is constant around $m^{uu} \sim 0.8$ under uncoordinated punishment but displays a substantial fall under coordinated punishment when $q = q'_2$. In particular, it starts by showing a steady pattern around $m^{cp} \sim 0.9$ and falls to $m \sim 0.5$ when this threshold is reached. From this point onward, m under coordinated punishment falls to $m \sim 0.2$. This result indicates that for a sufficiently high proportion of reciprocal investors, coordinated punishment also performs better in terms of the requirement of fair-minded allocators that sustain a pooling equilibrium.

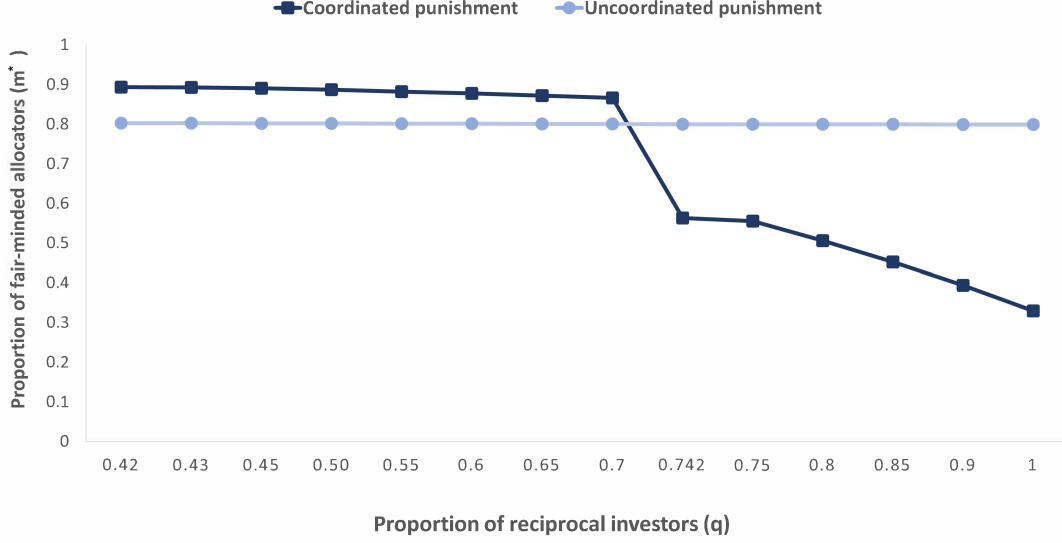


Figure 3: PBE fair-minded allocators of the numerical example for $q > q_1''$.

The superiority of the coordinated punishment scheme, which happens for a high proportion of reciprocators, is driven by the fact that reciprocators are willing to punish higher offers under coordinated punishment, to which selfish allocators return higher rewards than under uncoordinated punishment in order to avoid punishment. This, in turn, enhances efficiency through joint investment even with a lower proportion of fair-minded allocators in the population.

This effect is indeed given by coordination as a mechanism to overcome the free-rider issue. Intuitively, differences between the two types of punishment lie in the free-riding incentives between reciprocators. With coordinated punishment, this behaviour can be avoided when there are many reciprocators in the population as if one of the two investors free rides, punishment actions will be costly and ineffective in destroying the allocator's payoff. When the proportion of reciprocators is high, the probability of being paired with another reciprocator is also high, so the chance of there being joint punishment is higher. Uncoordinated punishment, however, is prone to free riding for any value of q , as the allocator's payoff is going to be undermined even if there is only one punisher. Hence, even if the proportion of reciprocators is high, the threat of joint punishment is not as strong as with coordinated punishment.

The superiority of coordinated punishment only with a high number of reciprocators is already present when considering that coordinated punishment has the same impact as uncoordinated punishment ($\lambda_c = 2\lambda_u$) as we do in this work. If, however, we considered

coordinated punishment to be strictly more destructive than uncoordinated punishment, following Calabuig and Olcina (2015), i.e. $\lambda_c > 2\lambda_u$, the differences would be even greater and this result would be even stronger.

4 Conclusions

In the theoretical comparison of uncoordinated and coordinated punishment, we have highlighted the superiority of the coordinated punishment scheme in achieving cooperation only when the proportion of reciprocal investors is sufficiently high.⁹ This occurs because there is a greater range of cases in which investors decide to jointly invest, the allocator returns a positive reward, and there is no punishment. Additionally, the amount returned under a coordinated punishment is larger than under uncoordinated punishment, and the proportion of fair-minded allocators required for this equilibrium to happen is lower.

The underlying reason for this behaviour is that the punishment threat exerted by coordinated punishers is stronger than the one exerted by uncoordinated punishers, as under a coordinated scheme, two investors willing to punish are needed for punishment to inflict damage on the allocators' payoffs. This is reflected in the fact that reciprocal investors are more demanding with the allocators' returns under a coordinated punishment system. As allocators anticipate this exigence, they will return rewards that are sufficiently high for the investors not to punish them, which enhances joint investment for a wider range of cases.

However, what is particularly noteworthy is that the superiority of coordinated punishment only occurs when there is a high proportion of reciprocators in the population. The reason behind this is that the coordination mechanism can avoid free-riding behaviour among reciprocators if many. With a low proportion of reciprocators in the population, however, such an advantage disappears.

Notice that the driving force of reciprocators punishing is the reduction of disadvantageous inequality, which can be understood as a public good achieved with such punishment actions. Nevertheless, the value that the two types of investors give to this public good

⁹Our result seems to support a possible co-evolution between institutions (coordinated punishment) and preferences for reciprocity in the investor population.

is not the same: only reciprocators value inequality reduction positively. A way of overcoming the inefficient outcome resulting from a public goods game is to make it become a coordination game, and, for this reason, coordinated punishment proves to be more efficient than uncoordinated punishment when there is a high proportion of reciprocators in the population. With a lower proportion of these, however, the presence of selfish investors dissipates the public good. In this case, the advantage of individual effective uncoordinated punishment prevails and differences between coordinated and uncoordinated punishment are minimal.

Appendix A. Proofs

Proof of Lemma 1. BNE of the punishment stage under coordinated punishment.

Let us first prove that selfish investors will never punish in the coordinated punishment stage by contradiction. For a selfish investor to implement punishment, the expected utility of punishing, $\mu \cdot U(p, p) + (1 - \mu) \cdot U(p, np)$, must be greater than the utility of not punishing, $\mu \cdot U(np, p) + (1 - \mu) \cdot U(np, np)$. By substituting the corresponding payoffs, the comparison the selfish investor makes is:

$$\mu(b - k - z/2) + (1 - \mu)(b - k - z) \geq b - k$$

Solving for μ we obtain that a selfish investor will punish if $\mu \geq 2$, which never holds.

Knowing a selfish investor is always going to choose np , a reciprocal investor can incorporate this in his/her comparison of the expected utility of punishing, $\mu \cdot U(p, p) + (1 - \mu) \cdot U(p, np)$ and his/her expected utility of not punishing $\mu \cdot U(np, p) + (1 - \mu) \cdot U(np, np)$. By substituting the corresponding payoffs, the comparison the reciprocal investor makes is:

$$\mu[b - k - z/2 - 2 \max\{0, (1 - b)(1 - \lambda_c) - b\}] + (1 - \mu)[b - k - z - 2 \max\{0, (1 - b) - b\}] \geq b - k - 2 \max\{0, (1 - b) - b\}$$

Notice that there are two types of inequalities in this comparison, but whether they become active or not depends on the value b takes. In particular, the inequalities $\max\{0, (1 - b) - b\}$ always become active for unfair offers of $b < 1/2$. If $b \leq b_1^c$, where $b_1^c = \frac{1 - \lambda_c}{2 - \lambda_c}$, inequality $\max\{0, (1 - b)(1 - \lambda_c) - b\}$ also becomes active.

Let us start with the case where $b \leq b_1^c$, such that all inequalities become active. Solving for μ we obtain that a reciprocal investor will punish if $\mu > \frac{z}{z/2 + 2\lambda_c(1 - b)}$, from which we define $\bar{\mu}$. This threshold $\bar{\mu}$ is always positive for $z \geq 0$, but for it to be $\bar{\mu} \leq 1$, then $b \leq b' = 1 - \frac{z}{4\lambda_c}$. As $b' > 1/2$, all values $b \leq b_1^c$ accomplish $\bar{\mu} \leq 1$. Therefore, in the reward segment $b \leq b_1^c$, if $\mu \leq \bar{\mu}$ then reciprocals would not punish, (np, np) , because the probability of being paired with another reciprocal investor is low. Otherwise, if $\mu > \bar{\mu}$, reciprocal investors would punish, (p, np) .

For the case where $b_1^c < b < 1/2$, only inequalities $\max\{(1-b) - b\}$ become active. Solving for μ we obtain that a reciprocal investor will punish if $\mu > \frac{z}{z/2+2(1-2b)}$, from which we define $\bar{\mu}$. This threshold $\bar{\mu}$ is always positive for $z \geq 0$, but for it to be $\bar{\mu} < 1$, then $b \leq b'' = \frac{1}{2} - \frac{z}{8}$. It can be shown that $b_1^c < b'' < 1/2$, from which we define two segments. If $b_1^c \leq b < b''$ whether there is punishment or not depends on $\bar{\mu}$. If $\mu \leq \bar{\mu}$ then (np, np) . Otherwise, if $\mu > \bar{\mu}$, then (p, np) . For $b'' \leq b < 1/2$, $\bar{\mu} > 1$, so $(np, np) \forall \mu$.

■

Proof of Lemma 2. BNE of the punishment stage under uncoordinated punishment.

The proof of not punishing being a dominant strategy for selfish investors in the uncoordinated punishment stage follows from the proof for lemma 1.

For the case of the reciprocal investor, the comparison that he/she makes is:

$$\mu[b - k - z/2 - 2 \max\{0, (1-b)(1-2\lambda_u) - b\}] + (1-\mu)[b - k - z - 2 \max\{0, (1-b)(1-\lambda_u) - b\}] \geq \mu[b - k - 2 \max\{0, (1-b)(1-\lambda_u) - b\}] + (1-\mu)[b - k - 2 \max\{0, (1-b) - b\}]$$

Notice that there are three types of inequalities in this comparison and whether they become active or not depends on the value of the reward b :

- If $b < b_1^u = \frac{1-2\lambda_u}{2-2\lambda_u}$, then all inequalities become active.
- If $b_1^u \leq b < b_2^u = \frac{1-\lambda_u}{2-\lambda_u}$, then inequalities $\max\{(1-b)(1-\lambda_u) - b\}$ and $\max\{(1-b) - b\}$ become active.
- If $b_2^u \leq b < 1/2$, then only inequality $\max\{(1-b) - b\}$ becomes active.

If $b \leq b_1^u$, all inequalities become active. Solving for μ we obtain that a reciprocal investor will punish if $\mu > \frac{2z-4\lambda_u(1-b)}{z}$, from which we define $\tilde{\mu}$. For this threshold to be positive, then $b \geq \hat{b} = 1 - \frac{z}{2\lambda_u}$, which holds by assumption 2. For this threshold to be less than 1, $b \leq b^* = 1 - \frac{z}{4\lambda_u}$. It can be shown that $b^* > b_1^u$, therefore, all values below b_1^u satisfy $\tilde{\mu} \leq 1$. Hence, we have two segments: $0 \leq b < \hat{b}$ and $\hat{b} \leq b < b_1^u$. For the first case, as $b < \hat{b}$, then $\tilde{\mu} < 0$ and therefore the unique BNE is $(p, np) \forall \mu$. For the second case, where $\hat{b} \leq b < b_1^u$, whether there is punishment or not depends on the value of μ . If

$\mu \leq \tilde{\mu}$ then (np, np) . Otherwise, if $\mu > \tilde{\mu}$, then (p, np) .

If $b_1^u \leq b < b_2^u$, then inequalities $\max\{(1-b)(1-\lambda_u) - b\}$ and $\max\{(1-b) - b\}$ become active. Solving for μ we obtain that a reciprocal investor will punish if $\mu \geq \frac{z-2\lambda_u(1-b)}{z/2+2(1-2b)-4\lambda_u(1-b)}$, from which we define $\tilde{\mu}$. Whether $\tilde{\mu}$ is positive is no longer straightforward and depends on the sign of the numerator and the denominator. For the numerator to be positive $b \geq \hat{b} = 1 - \frac{z}{2\lambda_u}$, which holds by assumption 2. For the denominator to be positive $b \leq \hat{\hat{b}} = \frac{1-2\lambda_u}{2-2\lambda_u} + \frac{z}{8(1-\lambda_u)}$ must hold. Finally, for $\tilde{\mu} \leq 1$, then b must satisfy $b < b^{**} = \frac{1-\lambda_u}{2-\lambda_u} - \frac{z}{4(2-\lambda_u)}$. By assumption 2, we can claim that $b_1^u < b^{**} < \hat{\hat{b}} < b_2^u$, defining three segments. For any b in the segment $b_1^u \leq b < b^{**}$, both the numerator and the denominator of $\tilde{\mu}$ are positive and $\tilde{\mu} \leq 1$ holds. Hence, whether there is punishment or not depends on $\tilde{\mu}$: if $\mu \leq \tilde{\mu}$ then (np, np) . Otherwise, if $\mu > \tilde{\mu}$, then (p, np) . For any b in the segment $b^{**} \leq b < \hat{\hat{b}}$, both the numerator and the denominator of $\tilde{\mu}$ are positive, but $\tilde{\mu} > 1$. Therefore, in this case the BNE is $(np, np) \forall \mu$. Finally, for the last segment where $\hat{\hat{b}} \leq b < b_2^u$, the numerator is positive, but the denominator is negative as $b \geq \hat{\hat{b}}$, therefore, $\tilde{\mu}$ must be lower than a negative number, which never holds, so (np, np) is the unique BNE $\forall \mu$.

In the last place, let us study the higher segment of b , where $b \geq b_2^u$. Now, only the inequality $\max\{(1-b) - b\}$ becomes active. Solving for μ , we obtain that for the reciprocal investor to punish $\mu \geq \tilde{\tilde{\mu}} = \frac{z-2(1-2b)}{z/2-2(1-2b)}$. As in the previous case, the numerator is positive by assumption 2, but for the denominator to be positive as well, it is necessary that $b \geq \hat{\hat{b}}' = \frac{1}{2} - \frac{z}{8}$. It can be shown that $b_2^u \leq \hat{\hat{b}}' < 1/2$, such that we have two segments. In the first segment, $b_2^u \leq b < \hat{\hat{b}}'$, the numerator of $\tilde{\tilde{\mu}}$ is positive while the denominator is negative. Therefore, $\tilde{\tilde{\mu}}$ must be lower than a negative number, which never holds, so (np, np) is the unique BNE $\forall \mu$. Alternatively, for $\hat{\hat{b}}' \leq b < 1/2$, both the numerator and the denominator are positive, but $\tilde{\tilde{\mu}} > 1$, so $(np, np) \forall \mu$.

■

Proof of Lemma 3: PBE of the reward stage under coordinated punishment:

For the reward policy with coordinated punishment, recall that allocators anticipate that only reciprocal investors may punish and that if this happens it does so for rewards $0 \leq b < b''$. Whether reciprocal investors punish or not depends on the beliefs of the

proportion of reciprocators in the population. In this range of values, selfish allocators can either set $b_s^c = 0$ and make profits of $\pi_s^c(0) = 2(1 - \mu^2\lambda_c)$, or offer a minimal reward $b_s^c > 0$ to avoid punishment and make profits of $\pi_{pm}^c(b) = 2(1 - b)$. Notice that the critical values $\bar{\mu}(b)$ and $\bar{\bar{\mu}}(b)$ are defined by the range of b where they are relevant. It can be easily shown that the critical values of μ are increasing with b and that they can be ordered such that $\bar{\mu}(0) \leq \bar{\mu}(b_1^c) = \bar{\bar{\mu}}(b_1^c) \leq \bar{\bar{\mu}}(b'') = 1$.

If the beliefs are very low, that is, if $\mu < \bar{\mu}(0)$, then $\mu < \bar{\mu}(b) \forall b$. Given that a selfish allocator faces no risk of being punished with so low beliefs, he/she will maximise his profits by returning $b_s^c = 0$.

If, instead, $\bar{\mu}(0) \leq \mu < \bar{\mu}(b_1^c) = \bar{\bar{\mu}}(b_1^c)$ a selfish allocator can either return nothing or set a minimal reward of $b^c = \frac{\mu(z/2 + 2\lambda_c) - z}{2\mu\lambda_c}$ to avoid punishment. This minimal reward can be directly obtained from the threshold $\bar{\mu}$. In order to choose between these two options, a selfish allocator will compare the expected profits of such options. In particular, he/she will offer $b_s^c = 0$ if $-(2\lambda_c^2)\mu^3 + (2\lambda_c + z/2)\mu - z > 0$. The roots of this cubic equation depend on its discriminant. If the discriminant is positive, then there is a real root, which is negative and two complex roots. In the positive domain, the cubic equation will always be strictly negative and, therefore, setting $b_s^c = b^c$ is better. If the discriminant is zero, then all roots are real, one of them negative and the other two would be two equal and positive roots. In the positive domain, the cubic equation will always be lower or equal than zero and, therefore, setting $b_s^c = b^c$ is better. However, if the discriminant is negative, then there are three real and unequal roots, one of them negative. Let us denote the two positive roots as μ_1 and μ_2 . In this case, for any $\mu < \mu_1$ and $\mu > \mu_2$, the cubic equation will be strictly positive and the positive reward $b_s^c = b^c$ would be returned, whereas for any $\mu_1 \leq \mu \leq \mu_2$, the cubic equation will be lower or equal than and the reward will be $b_s^c = 0$.

Finally, if $\bar{\mu}(b_1^c) = \bar{\bar{\mu}}(b_1^c) \leq \mu \leq \bar{\bar{\mu}}(b'') = 1$, the selfish allocator can, as before, either return nothing or return a positive reward that avoids punishment, $b^{cc} = \frac{\mu(2+z/2)-z}{4\mu}$, which can be directly obtained from $\bar{\bar{\mu}}$ and it can be shown that it is below the fair return. A selfish allocator will return $b_s^c = 0$ if $-(4\lambda_c)\mu^3 + (2 + z/2)\mu - z > 0$. If the discriminant is positive or equal to zero, the positive reward $b_s^c = b^{cc}$ would always be offered, following the same reasoning as before. For the cases where the discriminant is negative, we denote as μ'_1 and μ'_2 the positive roots of the cubic equation. Hence, for any $\mu < \mu'_1$ and $\mu > \mu'_2$, the positive reward $b_s^c = b^{cc}$ would be returned, whereas for any $\mu'_1 \leq \mu \leq \mu'_2$, the reward would be $b_s^c = 0$.

■

Proof of Lemma 4: PBE of the reward stage under uncoordinated punishment:

For the reward policy with uncoordinated punishment, selfish allocators can either set $b_s^u = 0$ and make profits of $\pi_s^c(0) = 2(1 - 2\mu\lambda_u)$, or offer a minimal reward $b_s^u > 0$ to avoid punishment and make profits of $\pi_s^u(b) = 2(1 - b)$. Notice that the critical values $\tilde{\mu}(b)$ and $\tilde{\tilde{\mu}}(b)$ are defined by the range of b where they are relevant. It can be shown that $\tilde{\mu}(\hat{b}) \leq \tilde{\mu}(b_1^u) = \tilde{\tilde{\mu}}(b_1^u) \leq \tilde{\tilde{\mu}}(b^{**}) = 1$.

If the beliefs are very low, that is, if $\mu < \tilde{\mu}(\hat{b})$, then $\mu < \tilde{\mu}(b) \forall b$. Given that a profit maximiser faces no risk of being punished with so low beliefs, he/she will maximise his profits by returning $b_s^u = 0$.

If instead, $\tilde{\mu}(\hat{b}) \leq \mu < \tilde{\mu}(b_1^u) = \tilde{\tilde{\mu}}(b_1^u)$ a selfish can either return nothing or set a minimal reward of $b^u = \frac{\mu z/2 - z + 2\lambda_u}{2\lambda_u}$ to prevent punishment. This minimal reward can be directly obtained from the threshold $\tilde{\mu}$ and it can be shown that it is less than a fair return. In order to choose between these two rewards, a selfish allocator will compare the expected profits of such options. In particular, he/she will offer $b_s^u = 0$ if $\mu \leq \frac{2\lambda_u - z}{4\lambda_u^2 - z/2}$, which we denote as $\tilde{\mu}(b^u)$. Hence, for $\tilde{\mu}(\hat{b}) \leq \mu < \tilde{\mu}(b^u)$, the selfish allocator will return $b_s^u = 0$ and for $\tilde{\mu}(b^u) \leq \mu < \tilde{\mu}(b_1^u)$, the selfish allocator will return $b_s^u = b^u$.

Finally, if $\tilde{\mu}(b_1^u) = \tilde{\tilde{\mu}}(b_1^u) \leq \mu < \tilde{\tilde{\mu}}(b^{**}) = 1$, a selfish allocator can either return nothing or set a minimal reward of $b^{uu} = \frac{\mu(z/2 + 2(1 - 2\lambda_u)) - z + 2\lambda_u}{2(\lambda_u(1 + 2\mu))} + 2\mu$ to avoid punishment. Such b can be directly obtained from the threshold $\tilde{\tilde{\mu}}$ and it can be shown that it is less than a fair return. A selfish allocator will return $b_s^u = 0$ if $-\mu^2(8\lambda_u(1 + \lambda_u)) + \mu(z/2 + 2(1 - 2\lambda_u - 2\lambda_u^2)) - z + 2\lambda_u > 0$. The roots of this quadratic equation depend on its discriminant. If the discriminant is negative, then there are two complex roots, the quadratic equation will always be strictly negative and, therefore, the allocator will always return $b_s^u = b^{uu}$. If the discriminant is zero, then there is only one real root, the quadratic equation will always be lower or equal than zero and the selfish allocator will also return $b_s^u = b^{uu}$ always. Finally, if the discriminant is positive, then there are two real roots, one negative and one positive (μ_1''). In this case, for any $\mu < \mu_1''$, the zero reward $b_s^u = 0$ would be returned, whereas for any $\mu_1'' \leq \mu$, the positive reward would be $b_s^u = b^{uu}$.

■

Proof of Proposition 1:

If $\bar{q}(0) \leq q < q_1$ or $q_2 \leq q < \bar{q}(b_1^c)$, we have that fair-minded allocators (proportion m) will return the fair reward of $b_f^c = 1/2$, selfish allocators (proportion $1 - m$) will return $b_s^c = b^c$ and no type of investor will punish. For a selfish investor to invest in this case, it must hold that $m(0.5) + (1 - m)b^c - k \geq 0$, from which we obtain the critical threshold $m_1 = \frac{k-b^c}{0.5-b^c}$. On the other side, a reciprocal investor will invest if $m(0.5) + (1 - m)[b^c - 2(1 - 2b^c)] \geq 0$, from which we obtain the critical threshold $m^c = \frac{k-b^c+2(1-2b^c)}{0.5-b^c+2(1-2b^c)}$. Given that $m^c \geq m_1$, if $m \geq m^c$ holds, both types of investors will invest.

If $\bar{q}(b_1^c) \leq q \leq q_1'$ or $q_2' \leq q \leq 1$, we have that fair-minded allocators (proportion m) will return the fair reward of $b_f^u = 1/2$, selfish allocators (proportion $1 - m$) will return $b_s^c = b^{cc}$ and no type of investor will punish. The same reasoning applies for selfish and reciprocal investors to invest where the critical thresholds are $m_2 = \frac{k-b^{cc}}{0.5-b^{cc}}$ and $m^{cc} = \frac{k-b^{cc}+2(1-2b^{cc})}{0.5-b^{cc}+2(1-2b^{cc})}$ respectively. Given that $m^{cc} \geq m_2$, if $m \geq m^{cc}$ holds, both types of investors will invest.

■

Proof of Proposition 2:

If $\tilde{q}(b^u) \leq q < \tilde{q}(b_1^u)$, we have that fair-minded allocators (proportion m) will return the fair reward of $b = 1/2$, selfish allocators (proportion $1 - m$) will return $b_s^u = b^c$ and no type of investor will punish. For a selfish investor to invest in this case, it must hold that $m(0.5) + (1 - m)b^u - k \geq 0$, from which we obtain the critical threshold $m_3 = \frac{k-b^u}{0.5-b^u}$. On the other side, a reciprocal investor will invest if $m(0.5) + (1 - m)[b^u - 2(1 - 2b^u)] \geq 0$, from which we obtain the critical threshold $m^u = \frac{k-b^u+2(1-2b^u)}{0.5-b^u+2(1-2b^u)}$. Given that $m^u \geq m_3$, if $m \geq m^u$ holds, both types of investors will invest.

If $q_1'' \leq q \leq 1$, we have that fair-minded allocators (proportion m) will return the fair reward of $b = 1/2$, selfish allocators (proportion $1 - m$) will return $b_s^u = b^{uu}$ and no type of investor will punish. The same reasoning applies for selfish and reciprocal investors to invest where the critical thresholds are $m_4 = \frac{k-b^{uu}}{0.5-b^{uu}}$ and $m^{uu} = \frac{k-b^{uu}+2(1-2b^{uu})}{0.5-b^{uu}+2(1-2b^{uu})}$ respectively. Given that $m^{uu} \geq m_4$, if $m \geq m^{uu}$ holds, both types of investors will invest.

■

Proof of Proposition 3:

In the first place, notice that $\frac{\partial m^{cc}}{\partial b^{cc}}, \frac{\partial m^{uu}}{\partial b^{uu}} < 0$, which implies that m^{cc} and m^{uu} are both decreasing in b^{cc} and b^{uu} respectively. This implies that if $b^{cc} > b^{uu}$, then $m^{cc} < m^{uu}$.

In the next place, let us show when is $b^{cc} > b^{uu}$, where $b^{cc} = \frac{q(z/2+2)-z}{4q}$ and $b^{uu} = \frac{q(z/2+2(1-2\lambda_u))-z+2\lambda_u}{2(\lambda_u(1+2q))} + 2q$. Let us simplify these expressions by saying that $b^{cc} = \frac{A}{B}$ and $b^{uu} = \frac{A+C}{B+D}$, where $A = q(z/2+2) - z$, $B = 4q$, $C = 2\lambda_u - 4q\lambda_u$ and $D = 2\lambda_u(1+2q)$.

Notice that $D > 0$, but the sign of C depends on the value that q takes. In particular, if $q > 1/2$, then $C < 0$ and therefore $b^{cc} > b^{uu}$. However, if $q < 1/2$, then $C > 0$. In this case, $b^{cc} > b^{uu}$ if $\frac{A}{B} > \frac{C}{D}$, which holds when $8q^2 + (z/2 - 2)q - z > 0$. This quadratic equation has two real roots, one negative and one positive, being the positive root $q^* = \frac{(2+z/2)+\sqrt{(2-z/2)^2+32z}}{16}$. It can be proved that as $z \leq 2\lambda_u$, then $z \leq 2$, so $q^* < 1/2$. Therefore, for any $q > q^*$, $b^{cc} > b^{uu}$.

Finally, the highest positive root of the cubic equation $g(\mu)$ we consider in Proposition 1b) can be proved to be greater than this threshold, i.e. $q'_2 > q^*$. Thus, for any $q \geq q'_2$, $b^{cc} > b^{uu}$ always holds.

■

Proof of Proposition 4:

In order to prove the existence of an interval $q''_1 < q \leq q'_2$, where $b_s^u = b^{uu}$ and $b_s^c = 0$, we must show that $q''_1 < q'_2$ or that $q'_2 - q''_1 > 0$. These threshold are q'_2 , which is the highest positive root of $g(\mu)$ and q''_1 , which is the positive root of $h(\mu)$. These roots depend on the values of z and λ_c . It can be proved that, if we fix λ_c at every possible value satisfying Assumption 3 and evaluate $q'_2 - q''_1$ for every possible value of z satisfying Assumption 2, we find that $q'_2 - q''_1 > 0$. Similarly, if we fix z at every possible value satisfying Assumption 2 and evaluate $q'_2 - q''_1$ for every possible value of λ_c satisfying Assumption 3, we find that $q'_2 - q''_1 > 0$. Therefore, $q'_2 - q''_1 > 0$ always holds.

■

Appendix B. Other pooling PBE

Besides the PRPE, there is a set of pooling equilibria where, if the proportion of fair-minded allocators is remarkably high, selfish allocators can return nothing and still avoid punishment.

Proposition 5: Null return pooling equilibria with no punishment

- a) *In the team trust game with **coordinated punishment**, if $q < \bar{q}(0)$ and $m \geq m^0$, there exists a PBE in which both types of investors decide to invest, (I, I) , a fair-minded allocator returns the fair reward $b_f^c = 1/2$, a selfish allocator returns $b_s^c = 0$ and there is no punishment.*
- b) *In the team trust game with **uncoordinated punishment**, if $q < \tilde{q}(b_1^u)$ and $m \geq m^0$, there exists a PBE in which both types of investors decide to invest, (I, I) , a fair-minded allocator returns the fair reward $b_f^u = 1/2$, a selfish allocator returns $b_s^u = 0$ and there is no punishment.*

where: $m^0 = \frac{2+k}{2+0.5}$

For a sufficiently low proportion of reciprocal investors, selfish allocators anticipate no future punishment, following lemmas 1 and 2. In this line, they will maximise their payoffs by not returning anything at all, by lemmas 3 and 4. Even in this case there can be investment by both types of investors if the proportion of fair-minded allocators who return them a fair reward is high enough.

Finally, this game also has pooling equilibria which are inefficient, in the sense that reciprocators implement punishment given that selfish allocators return nothing and the proportion of fair-minded is not sufficiently high.

Proposition 6: Null return pooling equilibria with punishment

- a) *In the team trust game with **coordinated punishment**, if either $q_1 \leq q < q_2'$ or $q_1' \leq q < q_2'$, and $m \geq m^{cp}$, there exists a PBE in which both types of investors decide to invest, (I, I) , a fair-minded allocator returns the fair reward $b_f^c = 1/2$, a selfish allocator returns $b_s^c = 0$ and there is punishment from reciprocal investors.*

b) In the team trust game with **uncoordinated punishment**, if $\tilde{q}(b_1^u) \leq q < q_1''$ and $m \geq m^{up}$, there exists a PBE in which both types of investors decide to invest, (I, I) , a fair-minded allocator returns the fair reward $b_f^u = 1/2$, a selfish allocator returns $b_s^u = 0$ and there is punishment from reciprocal investors.

$$\text{where: } m^{cp} = \frac{2+k+z-q(z/2+2\lambda_c)}{2+0.5+z-q(z/2+2\lambda_c)} \text{ and } m^{up} = \frac{2(1-\lambda_u)+k+z-q(z/2+2\lambda_u)}{2(1-\lambda_u)+0.5+z-q(z/2+2\lambda_u)}.$$

If beliefs about reciprocal investors are sufficiently high such that there is conditional punishment from their part, but selfish allocators do not return anything to investors, there will be punishment by the reciprocal investors (lemmas 1 and 2). As before, even in this case, there could be incentives to invest in the project if the proportion of fair-minded allocators is large enough. This decision is now not only going to depend on the investment cost k , the reward b and the inequality measure $\alpha = 2$, but also on the punishment cost z and on the probability of being paired with another reciprocal investor, q with whom to share the cost. For the case of coordinated punishment, recall that the presence of another reciprocal is needed for punishment to have a negative impact on the allocator's final payoff.

References

- [1] Boyd, R., Gintis, H., & Bowles, S. (2010). “Coordinated punishment of defectors sustains cooperation and can proliferate when rare.” *Science*, 328(5978), 617-620.
- [2] Calabuig, V., Jiménez, N., Olcina, G., & Rodríguez-Lara, I. (2022). United we stand: On the benefits of coordinated punishment. ESI Working Paper 22-12. https://digitalcommons.chapman.edu/esi_working_papers/373/
- [3] Calabuig, V. & Olcina, G. (2015). “Coordinated punishment and the evolution of cooperation.” *Journal of Public Economic Theory*, 17(2), 147-173.
- [4] Diekmann, A., & Przepiorka, W. (2015). Punitive preferences, monetary incentives and tacit coordination in the punishment of defectors promote cooperation in humans. *Scientific Reports*, 5(1), 10321.
- [5] Fehr, E., & Gächter, S. (2000). “Fairness and retaliation: The economics of reciprocity.” *Journal of Economic Perspectives*, 14(3), 159-181
- [6] Fehr, E., & Schmidt, K. M. (1999). “A theory of fairness, competition, and cooperation.” *The Quarterly Journal of Economics*, 114(3), 817-868.
- [7] Gächter, S., Renner, E., & Sefton, M. (2008). “The long-run benefits of punishment.” *Science*, 322(5907), 1510-1510.
- [8] García, J., & Traulsen, A. (2019). Evolution of coordinated punishment to enforce cooperation from an unbiased strategy space. *Journal of the Royal Society Interface*, 16(156), 20190127.
- [9] Gürer, Ö., Irlenbusch, B., & Rockenbach, B. (2006). “The competitive advantage of sanctioning institutions.” *Science*, 312(5770), 108-111.
- [10] Molleman, L., Kölle, F., Starmer, C., & Gächter, S. (2019). People prefer coordinated punishment in cooperative interactions. *Nature Human Behaviour*, 3(11), 1145-1153.