

Economic Growth Centre
Economics, School of Social Sciences
Nanyang Technological University
14 Nanyang Drive
Singapore 637332

INTERGROUP CONFLICT AVERSION WEAKENS INTRAGROUP COOPERATION

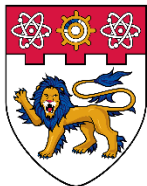
**Jonathan H.W. Tan
Friedel Bolle**

October 2019

EGC Report No: 2019/04

HSS-04-90A
Tel: +65 67906073
Email: D-EGC@ntu.edu.sg

Economic Growth Centre Working Paper Series



**NANYANG
TECHNOLOGICAL
UNIVERSITY**
SINGAPORE

Economic Growth Centre
School of Social Sciences

The author(s) bear sole responsibility for this paper.

Views expressed in this paper are those of the author(s) and not necessarily those of the
Economic Growth Centre, NTU.

Intergroup conflict aversion weakens intragroup cooperation*

Jonathan H W Tan,[†] Nanyang Technological University, Singapore

Friedel Bolle, European University Viadrina, Frankfurt (Oder), Germany

Sep 26, 2019

Abstract

We experimentally analyze dynamic strategies and social preferences in repeated intergroup conflicts where helping partners hurts rivals. Subjects are observed to be prosocial and conflict averse, in that they cooperate more when it benefits the ingroup but hold back when it inflicts losses on the outgroup. Our structural quantal response model with limited lookahead and group-specific altruism explains the data best: subjects expect partners to play cooperative strategies, whilst they are altruistic towards rivals. Social Value Orientation responses confirm that prosocials are relatively more conflict averse, which in turn weakens cooperation.

Keywords: cooperation; conflict; altruism; quantal response; experiment

JEL classification: C72, C91, D74, D91

* Acknowledgements: We thank Yves Breitmoser, Corentin Curchod, Simon Gächter, Ken Kamoche, Alexander Kritikos, George Kuk, Rajesh Kumar, Lim Wee Kiat, Daniele Nosenzo, David Saal, Steve Thompson, Yohanes Eko Riyanto, Zhao Zichen, and participants of workshops in Ningbo and Berlin for the advice and encouragement. We are grateful to Daniele Nosenzo, Shahin Bhugaloo, and Ruslan Kabalin for the experimental assistance, to CeDex and CRIBS for the logistical support, as well as to Nottingham University Business School for financial support under the Spark Fund.

[†] Corresponding author: Department of Economics, School of Social Sciences, Nanyang Technological University, 48 Nanyang Avenue, Singapore 639818, tel: 65-6790-4761, email: j.tan@ntu.edu.sg.

1. Introduction

Much of human cooperation in groups is set in the context of conflict between groups. In markets, conflicts arise between networks (e.g. competing airline alliances; Gimeno, 2004), between firms (e.g. law partnerships vying for a client; Leibowitz and Tollison, 1980), and within firms (e.g. between R&D teams; Birkinshaw, 2001). From a broader perspective of human evolution, groups fight for resources, dominance, and survival (Bernhard et al., 2006; Choi and Bowles, 2007; Norenzayan et al., 2016). In such conflicts, individuals face a dilemma between cooperating with partners for mutual gain, and defecting on partners to avoid hurting rivals. Ultimately, one must choose who to help and, conversely, who to hurt. We analyze how strategic reasoning, social preferences and social values drive behavior in such settings.

Cooperation (conflict) involves individually costly actions that increase (decrease) another's payoff via *synergy* (*rivalry*) between partners (rivals). Such incentive structures can strengthen group identification (Tan and Zizzo, 2008), which operates through group-specific preferences to increase cooperation, e.g. more charity to and less envy of the ingroup (Chen and Li, 2009). However, it can also be socially inefficient, e.g. turf wars within a conglomerate (Birkinshaw, 2001) or medical group (Gaynor, 1989). Parochial acts of altruism that benefit the ingroup at the cost of the outgroup are essential to the evolution of societies with repeatedly interacting groups (Choi and Bowles, 2007). In turn, evolutionary dynamics are shaped by strategic anticipation of and responses to others' behavior.

Our paper departs from previous research on social preferences and identity (Chen and Chen, 2011; Chen and Li, 2009), which focused on one-shot games where one's interactions with the ingroup and outgroup are independent. Our work is different in two ways. First, the game is repeated, to observe dynamic strategies in anticipation of and response to the strategies of partners and rivals. Second, group interests (i.e. welfare) are in direct conflict, to let us to observe preferences for partners and rivals. We systematically vary the structure of information (Fallucchi et al., 2013) and incentives (Reuben and Tyran, 2010) across games. Varying group-specific payoffs and social welfare implications while keeping the individual Nash equilibrium constant lets us disentangle prosociality and parochialism, thus differentiating it also from

previous research on parochial altruism (Aaldering et al., 2012; Abbink et al., 2012). This data allows us to experimentally and structurally estimate group-specific strategies and preferences.

We model intergroup conflict with a simple game in which individually costly actions result in ingroup gains and outgroup losses.¹ To observe preferences for the ingroup vis-à-vis outgroup without confounding efficiency concerns, we use a novel game where social welfare stays constant even when partners gain and rivals lose. A preference for the ingroup indicates parochialism. To test prosociality effects, keeping parochial effects constant, we compare this to a game where cooperation yields social welfare gains. To restrict parochialism, we have two further games where ingroup cooperation does not affect the outgroup. For direct comparability, in one of these games, subjects receive information on rivals' past actions to control for potential feedback effects germane to group conflict (Bornstein et al., 2002; Tan and Bolle, 2007).²

We observe that intergroup conflict is self-limiting. Relative to the game with payoff-irrelevant information feedback on rival effort, partner cooperation decreases if rivalry inflicts social loss. “Outgroup love” happens when rivals mutually acquiesce, which reduces partner cooperation. This negative effect dominates the positive “ingroup love” effect, while increasing social gains from cooperation attenuates it.

To estimate preferences and strategies, we focus on two classes of structural models. The first combines *exogenous beliefs* that depend on previous actions with altruistic preferences or inequality aversion. *Logit quantal responses* (McKelvey and Palfrey, 1995) to expected payoffs determine the probabilities of actions. The second concerns reciprocity and *endogenous beliefs*, which are determined as an *agent quantal response equilibrium* (McKelvey and Palfrey, 1995, 1998). Both classes of models are tested with or without *limited lookahead*, which anticipates payoffs also in the next period (Johnson et al., 2002).

¹ This simplification, e.g. relative to Tullock contests where relative effort affects winning probabilities (Sheremeta, 2018), limits confounds from cognitive overload (De Dreu et al., 2015).

² In general, relative group efforts determine competitive outcomes. Information feedback can operate through enhanced group salience and responses to rivalrous acts to increase group identity and ingroup love or outgroup hate (Tan and Zizzo, 2008), or simply serve as a motivational benchmarking device.

The combination of exogenous beliefs and limited lookahead best explains partner cooperation, which is fostered by the belief that partners play Grim strategies, whereas conflict aversion is captured by the structural estimates of altruism to the outgroup. This suggests that ingroup altruism is internalized in dynamic cooperative strategies working toward mutual benefit. To validate motives underlying the heterogeneity often observed in intergroup conflict (Sheremeta, 2018), we jointly analyze our behavioral data with self-reported data on individual differences in *Social Value Orientation* (Messick and McClintock, 1968). Subjects classified as prosocials are indeed significantly more conflict averse.

Section 2 reviews the relevance of preferences, reasoning, and individual differences to intergroup conflict. Section 3 presents the experimental games, logistics and procedure. Section 4 reports the behavioral patterns, structural model, and individual differences. Section 5 discusses our findings and concludes.

2. Background

2.1 Social Preferences

Parochial altruism can drive conflict between groups in contests where effort is exerted to increase a group’s probability of winning a prize, and is socially inefficient (Abbink et al., 2012). However, Reuben and Tyran (2010) find increased cooperation when competing groups can all win if they tie, i.e. without negative externalities for each other. Prosociality can even motivate personal sacrifice in the interest of the outgroup (Aaldering et al., 2013). Halevy et al. (2008) show that subjects prefer to help the ingroup without hurting the outgroup. This option increases welfare, and confounds group-specific prosociality.

To disentangle group-specific effects we shall adapt social preference models with separate weights for the ingroup and outgroup. The simplest model is linear *altruism (spite)* utilities that increase (decrease) in the payoffs of self and other, and implies a *positive (negative)* dependency on others payoffs, implying “ingroup love” (“outgroup hate”). Parochial altruism describes a preference for benefitting the ingroup even at a cost to the outgroup.

$$u^{Alt}(\mathbf{x}) = x_i + \alpha x_j$$

Inequity aversion (FS; Fehr and Schmidt, 1999) and *reciprocity* (CR; Charness and Rabin, 2002) are powerful theories that organize behavior across a variety of games. Inequity averse utility decreases with the payoff differences between self and others, and this effect is greater if one is envious when $x_j > x_i$ than if one is guilty $x_j < x_i$, i.e. $\alpha > \beta$. CR reciprocity additionally penalizes a co-player for making an inefficient move, which switches on the weight θ by setting $q_j = 1$, thereby decreasing (increasing) charity (envy) towards to this co-player.

$$u_i^{FS}(\mathbb{x}) = x_i - \sum_{j \neq i} \alpha (x_j - x_i) \cdot I_{x_i < x_j} + \sum_{j \neq i} \beta (x_j - x_i) \cdot I_{x_i \geq x_j} \quad (1)$$

$$u_i^{CR}(\mathbb{x}) = x_i - \sum_{j \neq i} (\alpha + \theta q_j)(x_j - x_i) \cdot I_{x_i < x_j} + \sum_{j \neq i} (\beta - \theta q_j)(x_j - x_i) \cdot I_{x_i \geq x_j} \quad (2)$$

Altruism can also be modeled in the sense of *reference dependent altruism* (RDA; Breitmoser and Tan, 2013; 2019), which is a good (out-of-sample) predictor for diverse static and dynamic games. It also explains behavior in demand bargaining and majority bargaining games, where observed behavior is incompatible with quasi-concave utility functions such as FS, CR, and CES altruism. RDA allows us to test the nature of conditionality on a reference point x' , which is one's ex ante expected payoff (*absolute* RDA) or on co-players' payoffs (*relative* RDA). Relative RDA allows for fairness concerns. The strength of altruism depends on whether one's payoff is less or more than the expected payoff or co-players' payoff, with a discontinuity at the point where payoff equals the reference point. We also extend RDA to consider the dynamic nature of our data by testing its conditionality on historical payoffs.

$$u_i^{RDA}(\mathbb{x}) = x_i + \sum_{j \neq i} \alpha x_j \cdot I_{x_i < x'} + \sum_{j \neq i} \beta x_j \cdot I_{x_i \geq x'} \quad (3)$$

2.2 Strategic Reasoning

The repeated nature of intergroup conflict in between social groups (Choi and Bowles, 2007; Norenzayan et al., 2016) or countervailing alliances (Gimeno 2004) involve anticipatory and responsive strategies that are absent in static settings. This is evident in finitely repeated contests, where subjects avoid conflict by coordinating through communication (Leibbrandt and Sääksvuori, 2012). Further, in a sequential move game where the first mover can mitigate eventual losses by striking first, Böhm et al. (2016) observe retaliatory and pre-emptive motives,

but find no evidence for spite-driven unilateral strikes. Amidst intergroup conflict, dynamic cooperative strategies are possibly also playing out between partners.

We analyze three main classes of dynamic strategies. First of all, one could follow *rule-based* strategies such as “tit-for-tat” (TFT) or “Grim” strategies used as responses to observed co-player behavior (Dal Bó and Fréchette, 2018). For example, Grim strategies are triggered off when a co-player defects, upon which one switches from cooperation to defection forever. Such rules have no explicit time horizon, and do not involve utility maximization. Implicitly, the belief of partners playing such strategies drives cooperation. Second, we could assume explicit exogenous belief formation based on the past behavior of co-players. Third, dynamics could be endogenously due to reciprocity or using past payoffs as reference points.³ We test these models under the assumption of bounded rationality.

Quantal responses relax the best response assumption by letting the profitability of an option determine its choice probability. The exogenous formation of beliefs is combined with logit quantal responses to these beliefs, and reciprocity is combined with the endogenous formation of beliefs by assuming logit *quantal response equilibrium* (QRE; McKelvey and Palfrey, 1995). Players expect others to mix strategies according to QRE, and expect others to expect the same. Further, players who operate with a two-period time horizon *limited lookahead* (LLA; Johnson et al., 2002) also expect their future selves (“agents”) to mix strategies, yielding an *agent QRE* (AQRE; McKelvey and Palfrey, 1998). Structural models of QRE allow the juxtaposing of social preferences and bounded rationality. We evaluate which explanation fits the data best, and to estimate preference parameters to compare the strengths of various motives.

2.3 Individual Differences

To validate motives and explore the heterogeneity commonly observed in intergroup conflict (Sheremeta, 2018), we jointly analyze behavioral data from games and self-reported data on individual differences in sociocognitive dispositions. We follow Offerman et al. (1996)

³ A fourth possibility is incomplete information about co-players’ preferences and updating beliefs after observing co-player behavior.

to classify subjects by their *Social Value Orientation* (SVO; Messick and McClintock, 1968), who show experimentally that *prosocials* (i.e. those who value their own and others' welfare) are more likely to cooperate than *individualists* (i.e. those who value their own welfare) or *competitive* types (i.e. those who maximize the difference in earnings between self and others).

Aaldering et al. (2012) show that when representing a group in negotiations against an adversary, SVO prosocials are more likely to sacrifice individual gain for group and social gain. Abbink et al. (2012) classified “prosocials” using a prisoner’s dilemma, and showed that they exerted more socially inefficient (i.e. wasteful) effort in contests. We posit that prosocials will respond positively to increasing the degree of gains from cooperation, and negatively to increasing the degree of losses from conflict. If so, this will support the link between the respective behaviors observed and social welfare concerns.

To further explore the sociocognitive effects on competitive behavior, we consider Regulatory Focus Theory (RFT; Higgins, 1997, 1998), which measures a subject’s focus on gain promotion and loss prevention. We posit that prevention (promotion) is sensitized to losses (gains) from rivalry (cooperation). We find that prevention focus, as measured by RFT, desensitizes one to potential gains from cooperation, but sensitizes one to potential and realized losses from conflict. Due to its exploratory nature we relegate further details on its postulated effects to Appendix A5.

3. The Experiment

3.1 Games

A game has four players $i = A, B, C, D$ who form two groups $J = \{A, B\}$ and $K = \{C, D\}$. The game is repeated for ten periods. In each period, a player has an endowment $\bar{e}$, which can be interpreted as a resource constraint. From $\bar{e}$, one can allocate effort to an individual project (Project I) or to a group project (Project II). We denote the effort by player i the group project as e_i . The total effort by group J is $e_J = e_A + e_B$, and $e_K = e_C + e_D$ for group K . Each unit allocated to the individual project is normalized to 1, and does not affect co-players’ payoffs. Cooperation in a group project has a positive impact on partners $r_{in}e_J$, where r_{in} captures the degree of

synergy between partners and $0.5 < r_{in} < 1$. Cooperation in the rival group has a negative impact $r_{out}e_K$, where r_{out} captures the degree of rivalry between groups and $r_{out} \leq 0$. Players are symmetric in $\bar{e}$, r_{in} , and r_{out} . The payoff that player i from group J makes is given by

$$x_i = \bar{e} - e_i + r_{in}e_J + r_{out}e_K. \quad (4)$$

This model captures potential gains from cooperation, while cooperation can be undermined by the fear of partner defection, and a group with a competitive advantage gains while its rivals while a competitive disadvantage suffer losses. If $r_{in} < 1$ and players are purely self-interested, they would prefer partners incur effort costs. When $r_{in} > 0.5$, payoffs accrue for both players if they cooperate, i.e. both put effort in the group project. When $r_{out} = 0$ ($r_{out} < 0$), groups are independent (rivalrous) because players do not (do) suffer losses due to the rivals' efforts.

Assuming players maximize individual payoffs, the Nash equilibrium (NE) is $c_i = 0$ for all players, i.e. no player ever invests in the group project because every such strategy is strictly dominated by zero effort. Alternatively, if players maximize group payoffs, i.e. each group acts as one player, then intragroup cooperation and intergroup conflict described by $c_i = e$ for all i is the unique NE. Finally, if all four players collude to maximize social welfare, then they choose to put effort in the individual project if $2(r_{in} + r_{out}) < 1$ and the group project if $2(r_{in} + r_{out}) > 1$. By backward induction of the finitely repeated game, the equilibrium of the stage game holds.

Our experiment has four games: *PD*, *INFO*, *RIVALRY*, and *SYNERGY*. Parameters satisfy the conditions required to maintain the unique individual NE of defection. For all games, $e = 150$ to avoid decimals in payoffs and to ease the calibration of marginal incentives. The *PD* is a control game where pairs of firms play in isolated groups ($r_{out} = 0$) and with synergy between partners ($r_{in} = 2/3$), but without information regarding what any other group has done. *INFO* differs from *PD* only in that we provide information on the actions of rivals after each period. This information is extraneous as it bears no pecuniary consequence. Observed responses to this information thus reflect concerns for non-pecuniary payoffs from comparison outcomes, using *PD* to control for the effect of pecuniary payoffs from partner actions.

RIVALRY is implemented by increasing r_{out} from 0 (in *INFO*) to $-1/6$. This manipulation allows us to test the effect of potential losses from intergroup rivalry, while controlling for the

effect of potential gains from intragroup cooperation by keeping it the same as in *PD* and *INFO*, and the effect of information by having the same feedback as in *INFO*. The value for rivalry is parameterized such that 1) the equilibrium strategy will be kept constant across games; 2) actions will not lead to efficiency losses overall, and; 3) the efficiency loss relative to the previous two games net out once synergy is increased.

SYNERGY introduces, relative to *RIVALRY*, an increase in r_{in} from $2/3$ to $5/6$, to test the effect of synergy in a competitive environment. The degree of losses inflicted on rivals is the same as that in *RIVALRY*. While this is no longer a zero-sum game, the game theoretic predictions are as before. This allows us to test the effect of potential gains from intragroup cooperation, while controlling for the effect of potential losses from intergroup conflict with *RIVALRY*. Also, the change in social welfare from cooperation is kept constant relative to the *PD* and *INFO*.

The social efficiency of cooperation in *PD*, *INFO*, and *SYNERGY* are equal, as each unit contributed to the group increases social welfare by 0.5. In *RIVALRY*, for each unit contributed, the group gains by the same amount as the rivals lose, and social welfare is kept constant across choices. This eliminates cooperative behavior motivated by efficiency concerns (Engelmann and Strobel, 2004). Its zero-sum nature has partners gain at the expense of self and rivals, allowing us to test one's parochial preference for the ingroup over the outgroup. In *SYNERGY*, each unit invested increases social welfare as the group's gain is double the loss of the rivals. *INFO* serves as the benchmark to control for feedback on others' behavior, which is essential to conflict. Such information can be used for social comparisons and to evaluate relative performance and efficacy in reaching goals (Festinger, 1954; Suls and Miller, 1977).

< Insert Table 1 about here.>

PD is presented in Table 1a and *RIVALRY* in Table 1b. Table 1c summarizes parameters and equilibrium predictions for individual, group, and society across games. This combination of games systematically isolates the effect of each treatment. In addition, the chosen parameter values stretch the variance of gains and losses, maintaining the unique individual equilibrium prediction across games, allowing for direct comparability and causal inference.

3.2 Logistics and Procedure

The experiment was conducted at the CeDEx experimental economics laboratory at the University of Nottingham. Subjects were recruited from a database of potential experimental participants using the recruitment software ORSEE (Greiner, 2004) and via email. They were undergraduates and postgraduates from various disciplines. The experiment was programmed in z-Tree (Fischbacher, 2007), and the profit calculator and questionnaire response interface (described below) were programmed in Visual Basic 6. The experimental instructions and control questionnaire were provided in print. See Appendix A1 for instructions and A2 for screenshots. Further experimental materials are available on request. We ran six sessions for each of the four games. Pairs of sessions of the same game were conducted simultaneously in one sitting to enhance anonymity. Each session had four subjects; each sitting had eight subjects. A total of 96 subjects participated. Games were counterbalanced in cycles and changed across each sitting. Subjects were randomly seated in the laboratory. Computer terminals were partitioned to prevent communication and seeing others' decisions and screens.

Pre-experimental instructions and control questionnaire. After being seated, subjects read the experimental instructions and answered a control questionnaire, to check their understanding of the rules. The experimental supervisor individually advised subjects with queries or with incorrect questionnaire answers. The text consistently employed a market frame, e.g. subjects were told that they were managers of firms belonging to an alliance and that they could earn profits by choosing between projects to invest in. This was to contextualize the task consistently with one of its main applications to make it more intuitive and engaging.

Training stage. This next stage allowed subjects to engage in an undirected use of the profit calculator for five minutes before playing the game. The training stage was primarily for ensuring that subjects had the chance to familiarize themselves with the payoff structure of the game, e.g. the consequences of different combinations of choices made by different co-participants. It can be interpreted as a form of pre-operational strategic evaluation stage, where partners, workers, or managers are able to familiarize themselves with different aspects of the

business setting. The profit calculator was displayed on a computer on each subject's right, and the game was played on a computer on each subject's left.

Game play. After the training phase, the computer program randomly assigned subjects to their respective groups. They were told their group assignments, and that the subjects they were matched with would remain constant throughout the experiment. Subjects did not know the group assignments of the other subjects in the room. They knew that there were two sessions running simultaneously, and that subjects for each session were randomly seated in the room. Subjects played ten periods of one of four games. We continued to provide the generic profit calculator, because it is a natural way for subjects to gather information about the consequences of alternative decisions.

After receiving feedback in each period, subjects then proceeded to the next period until all ten periods were done. In these games they could accumulate earnings in experimental points, each worth £0.01. The experiment was thus incentive compatible. To minimize noise from subject confusion, experimental points were scaled for cognitive simplicity and straightforward to convert to experimental earnings. This conversion rate was calibrated to yield an expected payment rate comparable with economics experiments ran in the UK.

Final questionnaire. After playing the experimental games, we collected self-reported data on age, gender, course of study, years of professional experience and occupation, and sociocognitive perspective of Social Value Orientation (Messick and McClintock, 1968) and Regulatory Focus (Higgins et al., 2001; see Appendix A5). We test SVO by asking subjects to choose between different hypothetical monetary allocations between self and other (following Offerman et al., 1996). It was noted that there were no right or wrong answers. Monetary incentives were not provided here. The inventory questions were presented and computational keys are as given by the original authors. Subjects entered their responses on the computer.

Payments. After completing the final questionnaire, the computer then reported to each subject how much one had earned, and this information was kept private and unknown to other subjects. Subjects were paid one at a time by the experimental supervisor. Payments were cumulative across periods. Subjects could earn a minimum of £10 and a maximum of £27.50

depending on game, their choices and those of their co-participants. Sessions lasted up to 1 hour. The average payment was approximately £17.

4. Results

To test treatment effects, we run univariate analyses comparing games using non-parametric Mann Whitney U tests. The first of two steps that we take to account for the potential non-independence of data due to feedback and interactions between subjects is to treat the average cooperation per session as an independent observation. We then run structural analyses of social preferences and bounded rationality. The results of our univariate tests shall also be corroborated by regression analysis. Multivariate analysis further allows us to test for interactions between subjects and responses to the behavior of others in this dynamic setting, but most importantly it lets us test for individual differences in sociocognitive effects. We use binary logit regressions with robust clustering at the session level, which is the second way by which we cater for possible non-independence across repeated interaction. Binary logit is used because of the nature of dependent variable, “cooperate” or “not cooperate.”

4.1 Behavioral Patterns

<Insert Figure 1 about here.>

Figure 1 presents the cooperation rates for each game, which are 0.39 in *PD*, 0.50 in *INFO*, 0.18 in *RIVALRY*, and 0.33 in *SYNERGY*. Figure 2 shows how cooperation rates change over time in each game. We use *INFO* as a benchmark that controls for the effect of information with feedback germane to conflict. The parameter values of *INFO* are the same as *PD*, but subjects in *INFO* also find out how much “rivals” are cooperating. Cooperation rates are not significantly higher in *INFO* ($z = 1.366$, $p = 0.172$). We test parochialism by comparing cooperation rates of *INFO* and *RIVALRY*, where r_{in} is equal and the feedback on the choices of partner and rival are present in both games, so the only difference is r_{out} in *RIVALRY*. For robustness we also compare *RIVALRY* with *PD* which has $r_{out} = 0$ and no feedback. Consistently, cooperation is significantly lower in *RIVALRY* than in *INFO* ($z = -2.169$, $p = 0.030$) and *PD* ($z =$

-2.242, 0.025). The *prosocial* effect of increased gains from cooperation r_{in} is tested by comparing the cooperation in *RIVALRY* and *SYNERGY*, because the only difference between the two games is positive net social benefit ($r_{in} + 2r_{out} > 0$) due to a higher r_{in} in the latter. The rate of cooperation is significantly higher in *SYNERGY* than in *RIVALRY* ($z = 2.082$, $p = 0.037$). To test the robustness of prosociality as a motive, we compare *INFO* and *SYNERGY* where actions yield similar welfare effects, and find no significant difference ($z = 0.724$, $p = 0.469$).

<Insert Table 2 and Figure 2 about here.>

Dynamics. Regression model 1 in Table 2 tests the robustness of the above treatment effects and dynamic effects due to experience and others' behavior. The dependent variable is *Coop* (= 1 if one chose Project II, i.e. to cooperate, or 0 if otherwise). The independent variables are *Period* (=1 to 10) to test for experience, lagged action variable *LPartnerCoop* (= partner's effort in previous period, i.e. 0 or 1) and *LRivalCoop* (= rivals' effort in previous period, i.e. 0, 1, or 2), to test for responses to partner and rival behavior in the previous period, respectively, and treatment variables namely the degree of synergy r_{in} (= 2/3 for *PD*, *INFO*, and *RIVALRY*, or = 5/6 for the *SYNERGY*) and the degree of rivalry r_{out} (= 0 for *PD* and *INFO*, or = 1/6 for the *RIVALRY* and *SYNERGY*). Model 1 analyzes data from all games and includes the treatment dummy *NoFeedback* (= 0 for *INFO*, *RIVALRY*, and *SYNERGY*, or = 1 for *PD*).

Cooperation decreases over time, as seen in Figure 2 and confirmed by the negative *Period* coefficient. *LPartnerCoop* is positive and significant indicating that subjects match the behavior of partners in previous periods. This shows that static theories cannot explain behavior in our experiment, which we analyze dynamically in Section 4.2. *Feedback*, which is inherent in *INFO*, *RIVALRY*, and *SYNERGY*, is not significant. *LRivalCoop* is not significant, explaining the lack of feedback effects. The positive and significant r_{in} coefficient confirms that cooperation increases with synergy. The negative and significant r_{out} coefficient confirms that cooperation decreases with rivalry. Our results on treatment effects are therefore robust.

4.2 Structural Analysis of Reasoning and Preferences

<Insert Figure 3 and Table 3 about here.>

Our dynamic analysis continues with Figure 3, which shows cooperation aggregated over games in each period as a function of own, partner, and rival actions in the previous period. Each line depicts the group's state in the previous period. From the leftmost to rightmost line, we have mutual defection, unilateral defection, unilateral cooperation, and mutual cooperation. Cooperation rates are highest after mutual cooperation, but do not vary as systematically with rivalry in the previous period. Table 3 shows the distribution of current actions disaggregated by game as a function of own and partner actions in the previous period (see Appendix A3 for conditionality also on rivals' actions). Defection ($LCoop = 0$) is unconditionally repeated, while cooperation ($LCoop = 1$) is conditionally repeated only if the partner has also cooperated ($LPartnerCoop = 1$).⁴ Fisher's tests in Table 3 show that $LCoop$ and $LPartnerCoop$ have strong overall impact ($p < 0.1$ or better). The pattern of decreasing cooperation over time shapes our structural model. We evaluate theories that predict behavior in period t based on behavior in $t - 1$. Our evaluation is therefore based on the frequencies of cooperating in 2×2 classes ($LCoop$ by $LPartnerCoop$) for PD , and in the $3 \times 2 \times 2 \times 3 = 36$ classes (game by $LCoop$, by $LPartnerCoop$, by $LRivalCoop$) for the other three games.

<Insert Table 4 about here.>

Table 4 summarizes the results of our structural analysis of various models and their fit based on the Akaike Information Criterion (AIC) and Bayes Information Criterion (BIC).⁵ We combine altruism or inequality aversion with exogenous beliefs (depending on previous period actions), or reciprocity that has changing coefficients of altruism or inequality aversion with endogenous beliefs (determined by QRE). Combining static utilities with endogenous beliefs predicts the same behavior in every period, which is obviously different from observed behavior as shown in Figure 3. Combining dynamic utility with exogenous beliefs introduces two competing sources of dynamic behavior. For static utilities with exogenous beliefs, we evaluate decisions from the viewpoint of player A, whose partner is player B and rivals are players C and D. We denote payoffs with x_i and utility with u_i .

⁴ Table A3 in the appendix shows that previous rival decisions have weak overall impact.

⁵ The Akaike Information Criterion is $AIC = 2d - 2\ln(LL)$ with number of parameters d . The Bayes Information Criterion is $BIC = \ln(O)*d - 2\ln(LL)$ with number of observations O (Schwarz, 1978).

For *PD* we assume that player A's utility in period t is

$$u_A(t) = x_A + \alpha_B * x_B, \quad \alpha_B \leq 1. \quad (5)$$

For the other three games, we assume

$$u_A(t) = x_A + \alpha_B * x_B + \alpha_{out} * (x_C + x_D), \quad \alpha_B, \alpha_{out} \leq 1. \quad (6)$$

Payoffs are expected payoffs determined by *exogenous beliefs* as follows.

- (i) If a player has cooperated in the previous period, she expects her partner to repeat his action.
- (ii) If a player has not cooperated in the previous period, she expects her partner to defect again with certainty or to cooperate again with probability η .
- (iii) If both rivals have contributed or both have not cooperated, they are expected to repeat their actions. If one of them has cooperated and the other not, the defecting player is expected to repeat his action and the cooperating player to repeat it with probability η .

For $\eta=1$, partners are expected to repeat their strategies from the previous period. For $\eta=0$, partners are expected to play Grim.

We first assume that players maximize $u_A(t)$ in every period. In the absolute RDA models, in *RDAexp* the reference payoff x'_i is the *expected payoff*, while *RDAhist* conditions on the *historical payoff*, i.e. one's previous period payoffs.⁶ The payoff x_A is always the expected payoff, depending on whether one cooperated ($Coop = 1$) or not ($Coop = 0$). Coefficients in the relative RDA models depend on payoffs relative to those of partners and rivals.

$$\alpha_B \text{ applies only if } x_A < x'_B, \text{ otherwise it is substituted by } \beta_B, \quad (7)$$

$$\alpha_{out} \text{ applies only if } x_A < \frac{x'_C + x'_D}{2}, \text{ otherwise it is substituted by } \alpha_B. \quad (8)$$

Finally, we extend the models by assuming that players have limited lookahead (*LLA*), i.e. they determine also the expected payoff $u_A(t + 1)$ of the next period, which depends on present period actions (shaping expectations for the next period) and optimal next period actions. In our models, looking ahead strengthens incentives for partner cooperation because one player's defection in period t implies high probabilities of both players defecting in $t + 1$ and therefore in all further periods. Letting players look ahead more than one period increases

⁶ Thus, *RDAhist* has a dynamic component. But as we want to compare it with *RDAexp* we evaluate this model under the same assumption on belief formation.

this incentive, but such an effect is not much different from an increased discount rate (which is allowed to be larger than 1). Players maximize

$$\mathbf{w}_A(t) = \mathbf{u}_A(t) + \delta * \mathbf{u}_A(t + 1). \quad (9)$$

For comparability with the other approaches, we include the final period of data in our estimations. For all models, we assume errors in the sense of QRE, such that the action that maximizes utility is not chosen with probability 1 but that the probability of cooperation is determined by logit quantal responses

$$p_A(\text{Coop} = 1) = \frac{e^{\lambda u_A(\text{Coop}=1)}}{e^{\lambda u_A(\text{Coop}=1)} + e^{\lambda u_A(\text{Coop}=0)}}. \quad (10)$$

For LLA models, u is substituted by w . Table 4 shows that exogenous beliefs outperforms endogenous beliefs, and that within the class of models with exogenous beliefs *Alt* outperforms *FS*. The model specifications for the other models (mixture of rules, dynamic utility and endogenous beliefs) are described in Appendix A3. Both versions of *FS* do not improve from extending the time horizon by one future period (*LLA*) while all altruism models improve significantly. The best-fitting models are *Alt-LLA* and *RDAhist-LLA*.

The absolute quality of each theory is evaluated by two measures. The first is a χ^2 test, which shows that the predicted frequencies are not significantly different from the empirical frequencies (see Table 5). The second is the best possible log-likelihood score, which is reached when a theory predicts exactly the empirical frequencies.⁷ This prediction implies $LL = -418.7$. The prediction of the same average frequency for all cases has $LL = -598.9$. Theories can improve on the latter, but can never reach the former. With $LL = -439.7$ for *RDAhist-LLA* and $LL = -441.4$ for *Alt-LLA*, the two models are closer to the upper bound than other theories.

<Insert Table 5 about here.>

The four best models are of altruism. Their parameter estimates are reported in Table 5 (the estimates for other models are reported in Appendix A3). The coefficients α_B , β_B , and η are not significantly different from zero. By and large, the objective function of the first three models in Table 2 can be written as

⁷ This is possible only as a population mixture of deterministic behavior. Actual behavior is different because players often respond differently to the same previous period pattern of decisions.

$$\mathbf{u}_A(t) = \mathbf{x}_A(t) + 0.6 * \mathbf{x}_{out}(t) \quad (11)$$

$$\mathbf{w}_A(t) = \mathbf{u}_A(t) + 0.55 * \mathbf{u}_A(t + 1). \quad (12)$$

A player's cooperation is fostered by her increased expected payoff in the next period, but hindered by her decreased expected payoff in the present period and by altruism towards rivals. Subjects expect partners' actions to be rule-based ($\eta=0$, i.e., Grim) and not utility-based, and this cooperative strategy accounts for the absence of altruism towards partners in (11).

4.3 Individual Differences in Sociocognitions

To understand the subject heterogeneity underlying these behavioral patterns, we extend regression Model 1 to analyze how treatments and sociocognitions are associated with dynamic competitive behavior in *INFO*, *RIVALRY*, and *SYNERGY*. Model 2 adds the lagged action variable *LRivalCoop* (= rivals' project choices in the previous period, i.e. 0 or 1) to test for responses to rivalrous behavior in the previous period, SVO types *Prosocial* (=1 if classified as prosocial by the SVO, or = 0 if individualist, competitive, or unclassifiable), and interactions of treatment dummies and lagged action variables by SVO types.

<Insert Figures 4 and 5 about here.>

We find that 38 of 96 subjects are prosocials, and 23 are unclassifiable. Model 2 shows that prosocials cooperate significantly more when $r_{out} = 0$, relative to when $r_{out} < 0$ and welfare losses result from rivalry ($r_{out} \times \text{Prosocial}$). Figure 4 illustrates this effect of reduced cooperation being more pronounced for prosocials as we move from *INFO* where $r_{out} = 0$ to *RIVALRY* and *SYNERGY* where $r_{out} < 0$. This confirms that prosociality strengthens intergroup conflict aversion and in turn weakens intragroup cooperation. That said, this effect is attenuated by prevention focus. Prevention focus induces intragroup cooperation in anticipation of or retaliation to group conflict (significant negative interaction effects with r_{out} and *LRivalCoop*, respectively). See Appendix A5 for the analysis of regulatory focus effects on conflict behavior.

5. Discussion

We analyzed the interplay of strategic reasoning and social preferences in situations where helping partners hurts rivals, and found strong evidence of conflict aversion. One faces a conflict of interest with partners and rivals, as cooperation can benefit the partner but also trigger off a war with rivals and hurt one in return. With intergroup conflict, intragroup cooperation is weakened because: 1) reduced incentives to pursuing cooperative strategies with partners, and; 2) a positive regard for social efficiency including the welfare of rivals.

Our structural analysis sheds light on strategic and motivational drivers of cooperation and conflict. These motivational patterns are explained by cooperation fostered by the belief of partners playing Grim and outgroup altruism. Cooperative strategies improve partner welfare and are therefore compatible with altruism towards partners. However, Grim strategies embody a mutual understanding to stay loyal to the common cause of increasing joint payoffs. This cooperation breaks down once someone cheats. In this sense, Grim strategies can embody a self-serving form of ingroup altruism in repeated games.

Outgroup altruism, though, motivates the avoidance of mutually destructive wars. An applicative example is conflict reduction through communication and coordination (Leibbrandt and Sääksvuori, 2012). Prosociality limits the willingness to be hostile to the outgroup especially when it does not sufficiently benefit the partner and self. We confirm that conflict aversion is most pronounced for prosocials who refrain relatively more from inflicting losses on rivals. They represent a broader sense of social interest.

Parochial altruism is an evolutionary force that benefits and protects the ingroup even if it entails aggression towards the outgroup (Bernhard et al., 2006; Choi and Bowles, 2007; Norenzayan et al., 2016). That said, our analysis shows how the conjunction of cooperative strategies and outgroup altruism drive cooperation within groups whilst motivating groups to refrain from mutual destruction, respectively. Lookahead reinforces mutual cooperation, but considerable instability in behavior arises from the expectation of partners playing Grim: one deviation from cooperation triggers non-cooperation in all further periods. These counteracting forces influence behavior in our four treatments with different reward structures.

Cooperation is in this sense fragile. Future work can design mechanisms to stabilize the joint pursuit of welfare improvements, or extend our approach to analyzing group-specific strategies and preferences in conflicts between endogenously formed groups (Herbst et al., 2015), managed teams (Eisenkopf, 2014), or naturally occurring groups (Chuah et al., 2016).

References

Aaldering, H., Greer, L.L., Van Kleef, G.A. and De Dreu, C.K., 2013. Interest (mis) alignments in representative negotiations: Do pro-social agents fuel or reduce inter-group conflict?. *Organizational Behavior and Human Decision Processes*, 120(2), pp.240-250.

Abbink, K., Brandts, J., Herrmann, B. and Orzen, H., 2012. Parochial altruism in inter-group conflicts. *Economics Letters*, 117(1), pp.45-48.

Akaike, H., 1974. A new look at the statistical model identification. *IEEE Transactions on Automatic Control*, 19 (6): 716–723.

Bernhard, H., Fischbacher, U. and Fehr, E., 2006. Parochial altruism in humans. *Nature*, 442(7105), p.912.

Birkinshaw, J., 2001. Strategies for managing internal competition. *California Management Review*, 44(1), pp.21-38.

Böhm, R., Rusch, H. and Gülerk, Ö., 2016. What makes people go to war? Defensive intentions motivate retaliatory and preemptive intergroup aggression. *Evolution and Human Behavior*, 37(1), pp.29-34.

Bornstein, G., Gneezy, U. and Nagel, R., 2002. The effect of intergroup competition on group coordination: An experimental study. *Games and Economic Behavior*, 41(1), pp.1-25.

Breitmoser, Y. and Tan, J.H.W., 2013. Reference dependent altruism in demand bargaining. *Journal of Economic Behavior & Organization*, 92, pp.127-140.

Breitmoser, Y. and Tan, J.H.W., 2019. Why should majority voting be unfair?. *Journal of Economic Behavior & Organization*, forthcoming.

Charness, G. and Rabin, M., 2002. Understanding social preferences with simple tests. *The Quarterly Journal of Economics*, 117(3), pp.817-869.

Chen, R. and Chen, Y., 2011. The potential of social identity for equilibrium selection. *American Economic Review*, 101(6), pp.2562-89.

Chen, Y. and Li, S.X., 2009. Group identity and social preferences. *American Economic Review*, 99(1), pp.431-57.

Choi, J.K. and Bowles, S., 2007. The coevolution of parochial altruism and war. *Science*, 318(5850), pp.636-640.

Chuah, S.H., Gächter, S., Hoffmann, R. and Tan, J.H.W., 2016. Religion, discrimination and trust across three cultures. *European Economic Review*, 90, pp.280-301.

- Crowe, E. and Higgins, E.T., 1997. Regulatory focus and strategic inclinations: Promotion and prevention in decision-making. *Organizational behavior and human decision processes*, 69(2), pp.117-132.
- Dal Bó, P. and Fréchette, G.R., 2018. On the determinants of cooperation in infinitely repeated games: A survey. *Journal of Economic Literature*, 56(1), pp.60-114.
- De Dreu, C.K., Dussel, D.B. and Velden, F.S.T., 2015. In intergroup conflict, self-sacrifice is stronger among pro-social individuals, and parochial altruism emerges especially among cognitively taxed individuals. *Frontiers in Psychology*, 6, p.572.
- Eisenkopf, G., 2014. The impact of management incentives in intergroup contests. *European Economic Review*, 67, pp.42-61.
- Engelmann, D. and Strobel, M., 2006. Inequality aversion, efficiency, and maximin preferences in simple distribution experiments: Reply. *American Economic Review*, 96(5), pp.1918-1923.
- Fallucchi, F., Renner, E. and Sefton, M., 2013. Information feedback and contest structure in rent-seeking games. *European Economic Review*, 64, pp.223-240.
- Fehr, E. and Schmidt, K.M., 1999. A theory of fairness, competition, and cooperation. *The Quarterly Journal of Economics*, 114(3), pp.817-868.
- Festinger, L., 1954. A theory of social comparison processes. *Human Relations*, 7(2), pp.117-140.
- Fischbacher, U., 2007. z-Tree: Zurich toolbox for ready-made economic experiments. *Experimental Economics*, 10(2), pp.171-178.
- Freitas, A.L., Liberman, N., Salovey, P. and Higgins, E.T., 2002. When to begin? Regulatory focus and initiating goal pursuit. *Personality and Social Psychology Bulletin*, 28(1), pp.121-130.
- Gaynor, M., 1989. Competition within the firm: theory plus some evidence from medical group practice. *The Rand Journal of Economics*, pp.59-76.
- Gimeno, J., 2004. Competition within and between networks: The contingent effect of competitive embeddedness on alliance formation. *Academy of Management Journal*, 47(6), pp.820-842.
- Greiner, B., 2004. The online recruitment system orsee 2.0-a guide for the organization of experiments in economics. *University of Cologne, Working Paper Series in Economics*, 10(23), pp.63-104.
- Halevy, N., Bornstein, G. and Sagiv, L., 2008. "In-group love" and "out-group hate" as motives for individual participation in intergroup conflict: A new game paradigm. *Psychological Science*, 19(4), pp.405-411.
- Herbst, L., Konrad, K.A. and Morath, F., 2015. Endogenous group formation in experimental contests. *European Economic Review*, 74, pp.163-189.
- Higgins, E.T., 1997. Beyond pleasure and pain. *American Psychologist*, 52(12), pp.1280-1300.
- Higgins, E.T., 1998. Promotion and prevention: Regulatory focus as a motivational principle. In *Advances in Experimental Social Psychology* (Vol. 30, pp. 1-46). Academic Press.

- Higgins, E.T., Friedman, R.S., Harlow, R.E., Idson, L.C., Ayduk, O.N. and Taylor, A., 2001. Achievement orientations from subjective histories of success: Promotion pride versus prevention pride. *European Journal of Social Psychology*, 31(1), pp.3-23.
- Johnson, E.J., Camerer, C., Sen, S. and Rymon, T., 2002. Detecting failures of backward induction: Monitoring information search in sequential bargaining. *Journal of Economic Theory*, 104(1), pp.16-47.
- Leibbrandt, A. and Sääksvuori, L., 2012. Communication in intergroup conflicts. *European Economic Review*, 56(6), pp.1136-1147.
- Leibowitz, A. and Tollison, R., 1980. Free riding, shirking, and team production in legal partnerships. *Economic Inquiry*, 18(3), pp.380-394.
- Marquardt, D.W., 1980. Comment: You should standardize the predictor variables in your regression models. *Journal of the American Statistical Association*, 75(369), pp.87-91.
- McKelvey, R.D. and Palfrey, T.R., 1995. Quantal response equilibria for normal form games. *Games and Economic Behavior*, 10(1), pp.6-38.
- McKelvey, R.D. and Palfrey, T.R., 1998. Quantal response equilibria for extensive form games. *Experimental Economics*, 1(1), pp.9-41.
- Messick, D.M. and McClintock, C.G., 1968. Motivational bases of choice in experimental games. *Journal of Experimental Social Psychology*, 4(1), pp.1-25.
- Norenzayan, A., Shariff, A.F., Gervais, W.M., Willard, A.K., McNamara, R.A., Slingerland, E. and Henrich, J., 2016. Parochial prosocial religions: Historical and contemporary evidence for a cultural evolutionary process. *Behavioral and Brain Sciences*, 39.
- Offerman, T., Sonnemans, J. and Schram, A., 1996. Value orientations, expectations and voluntary contributions in public goods. *The Economic Journal*, pp.817-845.
- Reuben, E. and Tyran, J.R., 2010. Everyone is a winner: Promoting cooperation through all-can-win intergroup competition. *European Journal of Political Economy*, 26(1), pp.25-35.
- Schwarz, G., 1978. Estimating the dimension of a model. *The Annals of Statistics*, 6(2), pp.461-464.
- Sheremeta, R.M., 2018. Behavior in group contests: A review of experimental research. *Journal of Economic Surveys*, 32(3), pp.683-704.
- Suls, J.M. and Miller, R.L., 1977. *Social comparison processes: Theoretical and empirical perspectives*. Hemisphere.
- Tan, J.H.W. and Bolle, F., 2007. Team competition and the public goods game. *Economics Letters*, 96(1), pp.133-139.
- Tan, J.H.W. and Zizzo, D.J., 2008. Groups, cooperation and conflict in games. *The Journal of Socio-Economics*, 37(1), pp.1-17.
- Zeelenberg, M., Van den Bos, K., Van Dijk, E. and Pieters, R., 2002. The inaction effect in the psychology of regret. *Journal of Personality and Social Psychology*, 82(3), p.314.

Tables and Figures

Figure 1. Mean cooperation by game

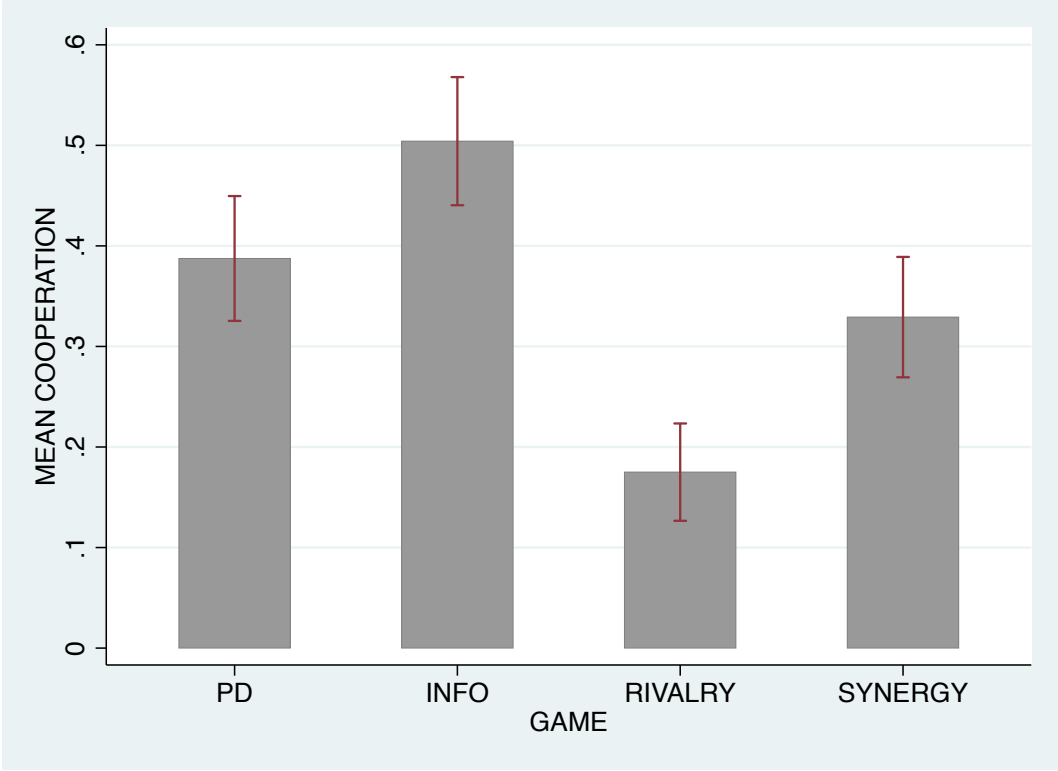


Figure 2. Mean cooperation over time

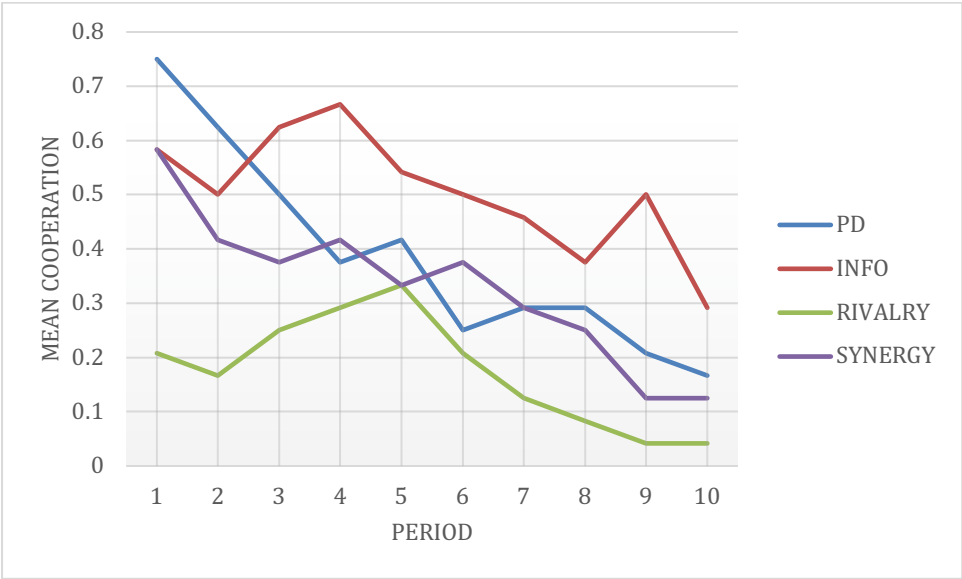


Figure 3. Cooperation as a function of own, partner, rival actions in the previous period

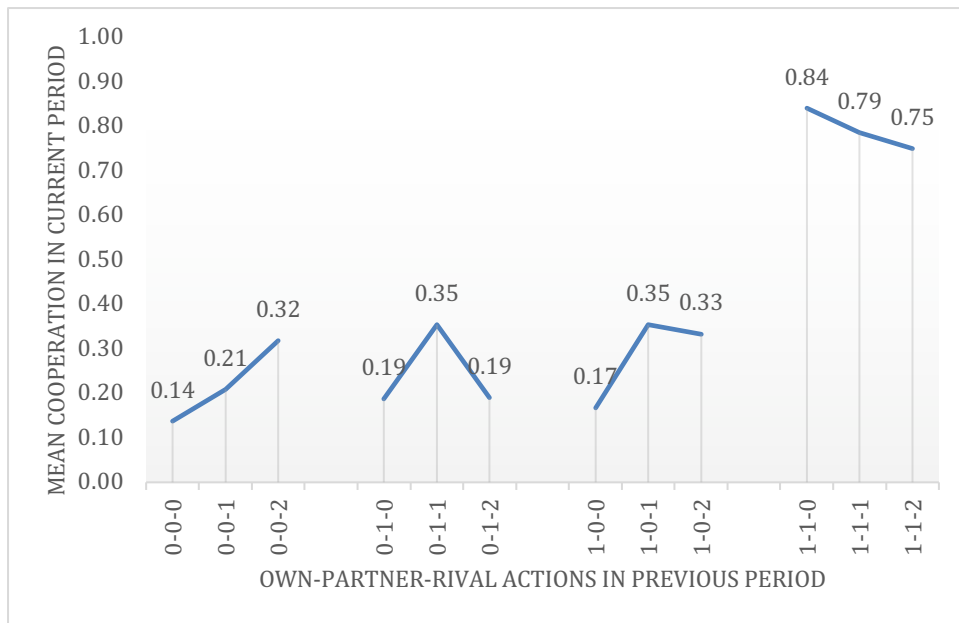


Figure 4. Cooperation as a function of game and SVO Prosociality

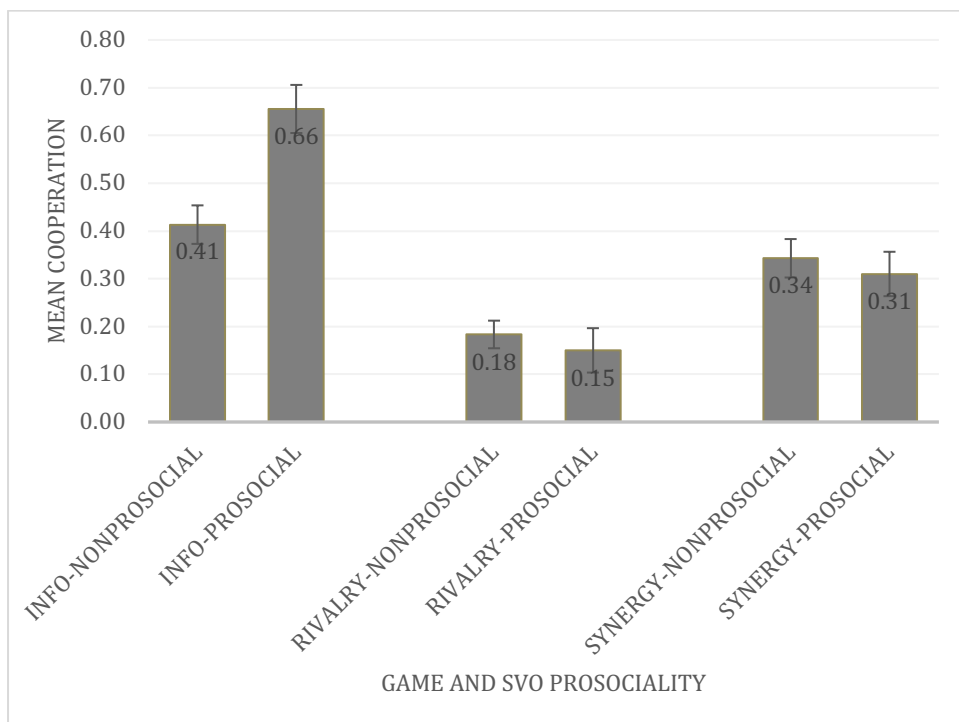


Table 1. Games, parameters, and predictions*

1a) Intragroup cooperation

	Cooperate	Defect
Cooperate	200, 200	100, 250
Defect	250, 100	150, 150

1b) Intragroup cooperation amidst intergroup conflict

	0 rival cooperates: $\pi=0$		1 rival cooperates: $\pi=25$		2 rivals cooperate: $\pi=50$	
	Cooperate	Defect	Cooperate	Defect	Cooperate	Defect
Cooperate	200, 200	100, 250	175, 175	75, 225	150, 150	50, 200
Defect	250, 100	150, 150	225, 75	125, 125	200, 50	100, 100

* Player 1 (2) receives the payoff stated on the left (right) of each cell, and chooses strategies by row (column). The outcome is determined by the intersection of both players' choices. In 1b, payoffs depend also on whether 0, 1, or 2 rivals cooperate within the other group, resulting in the payoff π from intragroup cooperation (given by the payoffs in Table 1a) minus losses 0, 25, or 50, respectively. The payoffs are for partners in the group, payoffs of rivals are determined respectively.

1c) Parameters, Nash equilibria, and social optima of the four games.

	<i>PD & INFO</i>	<i>RIVALRY</i>	<i>SYNERGY</i>
r_{in}	2/3	2/3	5/6
r_{out}	0	-1/6	-1/6
Individual NE	$e_i=0$	$e_i=0$	$e_i=0$
Group NE	$e_i=\bar{e}$	$e_i=\bar{e}$	$e_i=\bar{e}$
Social optimum	$e_i=\bar{e}$	indifferent	$e_i=\bar{e}$

Table 2. Binary logit regressions of cooperation*

	(1)	(2)
<i>Period</i>	-0.151*** (0.020)	-0.149*** (0.030)
<i>LPartnerCoop</i>	1.185*** (0.216)	0.989*** (0.377)
<i>NoFeedback</i>	-0.558 (0.411)	
<i>r_{in}</i>	3.438** (1.558)	3.588* (2.101)
<i>r_{out}</i>	-7.723*** (2.231)	-5.057* (3.041)
<i>LRivalCoop</i>		0.054 (0.270)
<i>Prosocial</i>		0.708 (2.078)
<i>r_{in} × Prosocial</i>		0.380 (3.496)
<i>r_{out} × Prosocial</i>		-7.699** (3.699)
<i>LPartnerCoop × Prosocial</i>		0.002 (0.327)
<i>LRivalCoop × Prosocial</i>		0.124 (0.613)
<i>Constant</i>	-2.034** (1.108)	-2.485* (1.417)
Pseudo <i>R</i> ²	0.1346	0.1510
<i>N</i>	864	648

* Robust standard errors in parentheses, clustered by session. The p-values reported here are two tailed, with * denoting $p < 0.10$, ** denoting $p < 0.05$, and *** denoting $p < 0.01$.

Table 3. Distribution of current actions as a function of own and partner actions in the previous period*

Condition	$LPartnerCoop$	$LCoop = 1$		$LCoop = 0$		Fisher's Test
		$Coop = 1$	$Coop = 0$	$Coop = 1$	$Coop = 0$	
<i>PD</i>	0	9	26	16	76	
<i>PD</i>	1	43	11	7	28	$<10^{-7}$
<i>INFO</i>	0	9	21	20	52	
<i>INFO</i>	1	74	10	4	26	0
<i>RIVALRY</i>	0	6	25	18	126	
<i>RIVALRY</i>	1	5	5	8	23	
<i>SYNERGY</i>	0	17	39	18	66	
<i>SYNERGY</i>	1	12	10	18	38	0.077

* $LCoop$ = lagged own cooperation, $LPartnerCoop$ = lagged partner cooperation, $Coop$ = current period cooperation. Two-sided Fisher Test between $d_{own} = Coop$ if $Coop = 1$ or 0.

Table 4. Overview of the performance of structural models*

Time horizon		Static: Present period				Limited Lookahead: Present and next period			
Characterization	Theory	#Par	-LogL	AIC	BIC	#Par	-LL	AIC	BIC
<i>Mixture of rules</i>	<i>Grim, ...</i>	7	473.6	961.1	994.5	-	-	-	-
<i>Static utility[§] & exogenous beliefs</i>	<i>Altruism</i>	4	525.4	1059.0	1071.0	5	441.4	892.8	916.6
	<i>FS</i>	5	464.9	939.8	944.5	6	453.0	918.0	943.0
	<i>RDAexp</i>	6	444.2	900.4	929.0	7	441.1	896.2	929.5
	<i>RDAhist</i>	6	455.9	923.8	952.4	7	439.7	893.4	926.7
<i>Dynamic Utility & endogenous beliefs</i>	<i>DynAlt</i>	5	517.7	1044.0	1069.0	6	480.6	973.2	1002.0
	<i>DynFS</i>	7	467.6	949.2	982.5	8	467.6	951.2	989.3

* BIC is computed with 864 data points. [§] *RDAhist* has a dynamic component but is estimated under the same assumptions as *RDAexp* in order to allow comparisons. *DynFS* = CR.

Table 5. Estimated parameters and their standard deviations for best altruism models*

	λ	β_B	α_B	β_{out}	α_{out}	η	δ	$p(\chi^2)$
<i>Alt-LLA</i>	6.930	-0.075		0.601		0.041	0.578	0.187
<i>S.D.</i>	1.435	0.204		0.170		0.089	0.072	
<i>RDAhist-LLA</i>	5.728	0.152	-0.005	0.705	0.732	0.099	0.543	0.268
<i>S.D.</i>	1.665	0.209	0.210	0.223	0.237	0.111	0.083	
<i>RDAexp-LLA</i>	9.69	-0.434	-0.217	0.417	0.422	0.017	0.551	0.135
<i>S.D.</i>	4.375	0.501	0.264	0.222	0.226	0.099	0.050	
<i>RDAexp</i>	3.735	1.000	0.602	0.697	0.724	0.055	-	0.044
<i>S.D.</i>	0.504	0.010	0.035	0.141	0.145	0.086	-	

* Evaluation of *Alt-LLA* ($LL = -441.4$, $AIC = 892.8$), *RDAhist-LLA* ($LL = -439.7$, $AIC = 893.4$), *RDAexp-LLA* ($LL = -441.2$, $AIC = 896.2$), and *RDAexp* ($LL = -444.2$, $AIC = 900.4$). Standard deviations (*S.D.*) are computed from the Hessian of the log-likelihood function. χ^2 is determined for the predictions of contribution frequencies in the 40 classes defined by the actions in the four treatments by self, partner, and rivals in the previous period.

Appendices

A1. Experimental Instructions

Instructions <PD>, [games with info and conflict]

Welcome to the experiment. There are 8 participants in this room. We are conducting two separate sessions in this room at the same time. Each session has 4 participants, randomly scattered across the room. Please do not talk to other participants, use your phone, look at others' screens, or read non experimental material. Your decisions are made under strict anonymity, and your earnings are kept strictly confidential, i.e. no other participant will be able to link them to your identity.

This experiment is made up of 2 stages. Before starting the experiment, you must first read these instructions, and then answer a questionnaire that will be used to check if you have understood the experimental tasks. When you have completed the questionnaire, please raise your hand and the laboratory assistant will come around to check your answers and whether you can proceed. Stage 2 contains the tasks you must perform. Stage 1 is a training stage to help you understand the task. We will first explain Stage 2.

Stage 2 is made up of 10 rounds. You will perform your tasks on the computer on your left. In each round, each participant will perform the task of making business decisions as the manager of a firm that they have been randomly assigned to. Firms A and B are members of Alliance X. Firms C and D are members of Alliance Y. The computer screen on your left tells you which firm you are managing and the alliance your firm is member of, and this will be specified in the top line of your computer screen throughout the experiment. The participant assigned to managing each firm will stay the same throughout the experiment. In each round, you will have the chance to earn profits in terms of points. Your experimental earnings will be based on the total number of points earned over the 10 rounds in Stage 2. Each point is worth £0.01.

Investments: A manager has to choose between investing in Project I or Project II.

- When a firm chooses Project I, it accrues R points, the other firm in the alliance accrues 0 points, and each firm in the other alliance accrues 0 points.
- For each firm in Alliance X that chooses Project II, Firm A accrues A_I points, Firm B accrues B_I points, Firm C accrues C_I points, and Firm D accrues D_I points.
- For each firm in Alliance Y that chooses Project II, Firm C accrues C_{II} points, Firm D accrues D_{II} points, Firm A accrues A_{II} points, and Firm B accrues B_{II} points.

A firm's *final profit* for a round is the total accrual of points resulting from the investment choices of both firms in its alliance [and both firms in the other alliance] in that round.

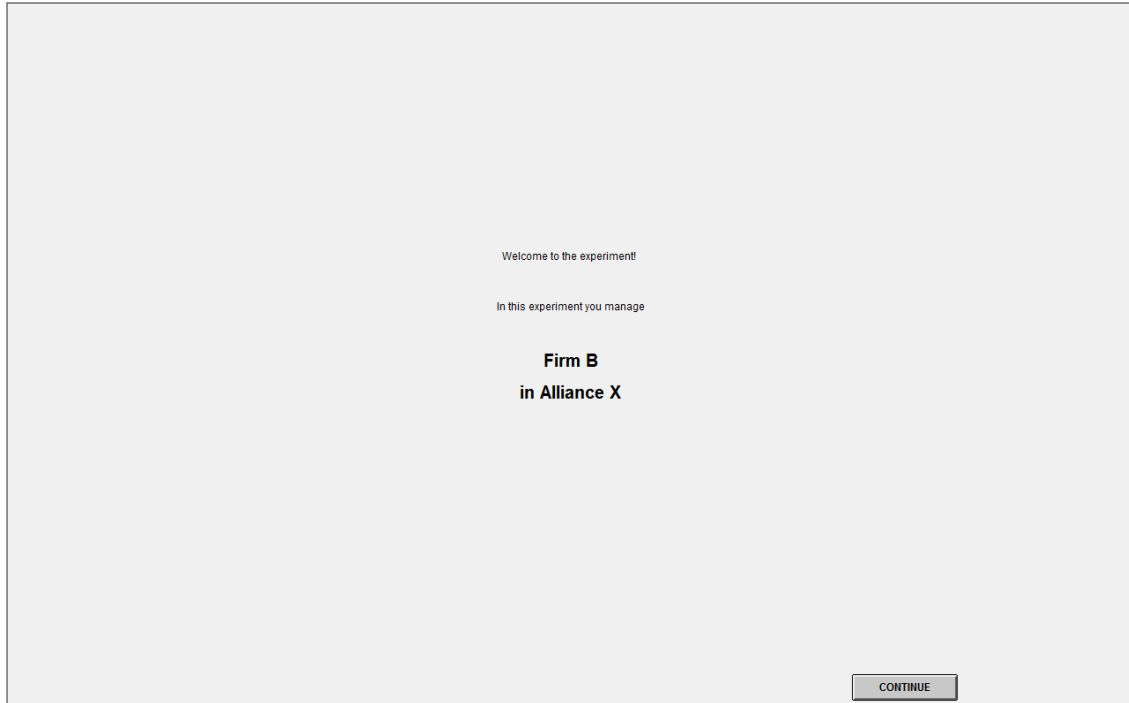
After all firms have made their investment choices, you will be informed of your final profit[,]
<and> the choices made in your alliance[, and the choices made in the other alliance] in that round. You will then proceed to the next round.

Stage 1 allows you to familiarize yourself with the task using a Profit Calculator, which is displayed on the computer screen on your right, before starting Stage 2. You can use it to calculate the final profit that each firm would accrue as a result of some combination of choices made by all firms. Please note that the “choices” on this calculator are hypothetical, and are only meant to aid your understanding of the task, and not for you to make choices on behalf of others. You have 5 minutes to use it before Stage 2 begins. The calculator will also be available for use throughout the experiment.

Please feel free to raise your hand if you have any queries at any point of time.

A2. Sample Screenshots of Experimental Program

First screen: welcome screen where subjects learn their firm and alliance. This screen stays on until subjects click CONTINUE or alternatively until 60sec have passed.



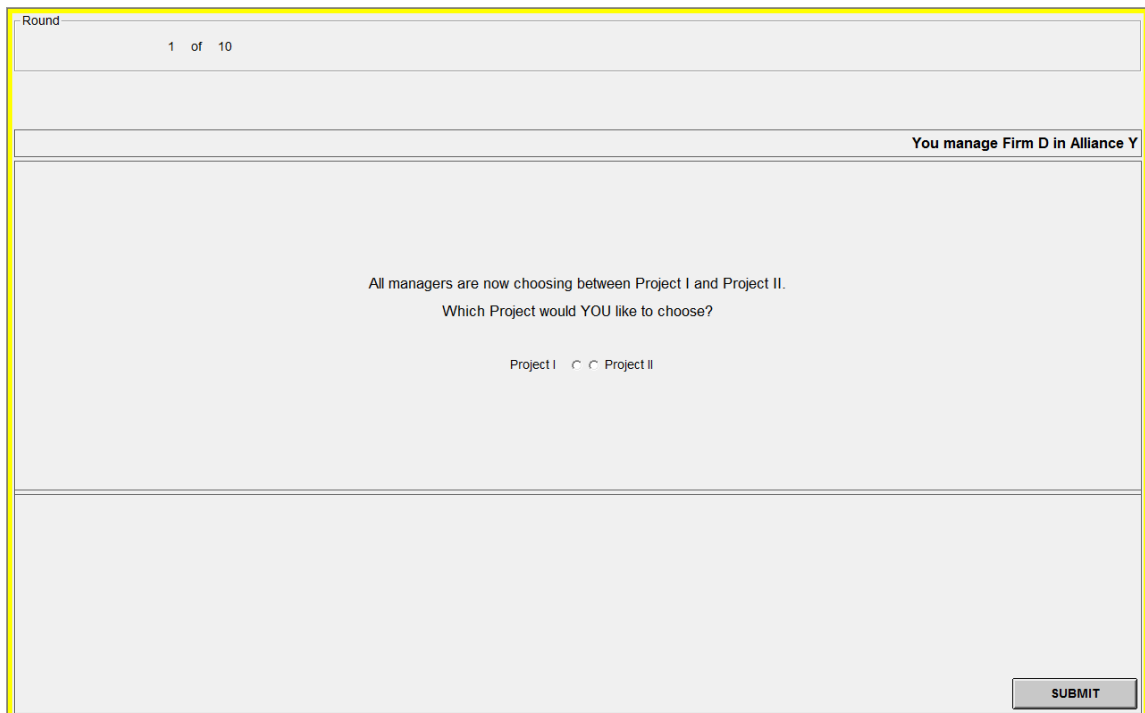
Welcome to the experiment!

In this experiment you manage

Firm B
in Alliance X

CONTINUE

Second screen: decision screen, it's the same for all subjects with the exception of the header that keeps track of their firm and alliance.



Round
1 of 10

You manage Firm D in Alliance Y

All managers are now choosing between Project I and Project II.
Which Project would YOU like to choose?

Project I ☐ Project II ☐

SUBMIT

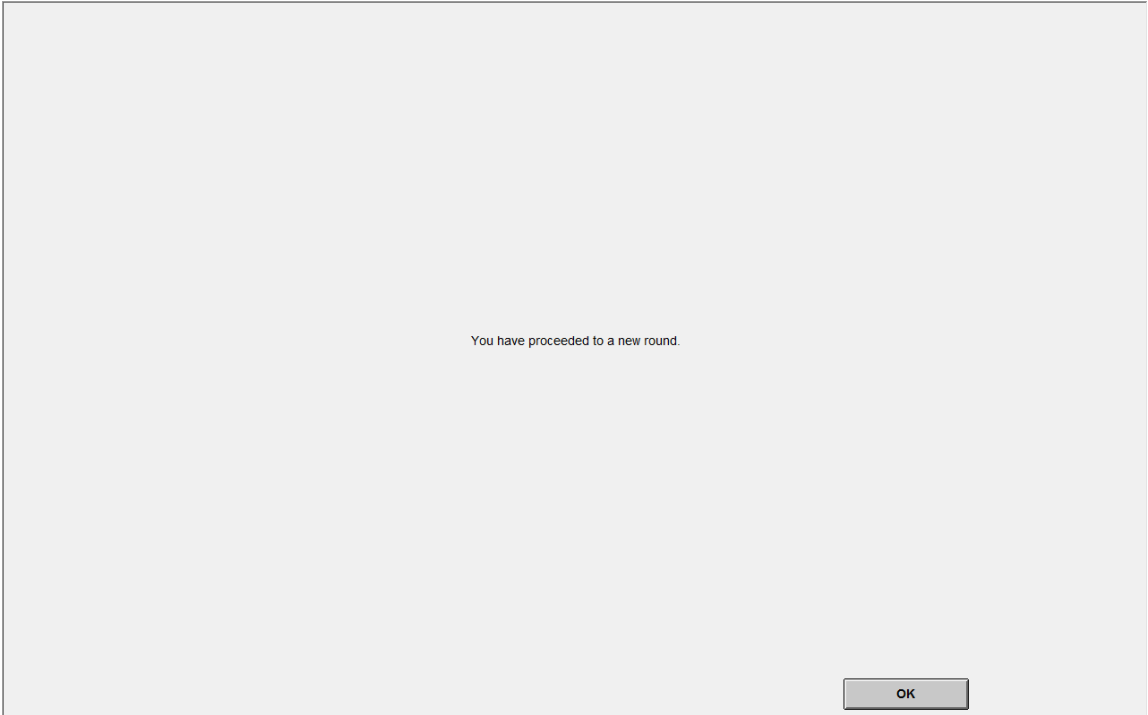
Subjects are asked to CONFIRM their decision and can still revise it before it becomes final (this is just for avoiding mistakes).

The screenshot shows a web-based decision task interface. At the top, it says "Round 1 of 10". Below this, a header bar indicates "You manage Firm B in Alliance X". The main area is titled "Decision Task". In the center, a dialog box is open with the following text: "You are about to submit your decision. To submit your decision, press CONFIRM. If you want to revise your decision, press REVISE." The dialog box has two buttons: "CONFIRM" and "REVISE". To the left of the dialog box, the text "All managers" is partially visible, and to the right, "d Project II." is visible. At the bottom right of the main area, there is a "SUBMIT" button.

Third screen: this screen shows the action/profit results for the current round. This info stays on screen until a subject clicks CONTINUE or alternatively until 120 sec have passed.

The screenshot shows the "Third screen" of the experiment. At the top, it says "Round 1 of 10". Below this, a header bar indicates "You manage Firm D in Alliance Y". The main area displays the following information: "In Alliance Y 1 Firm chose Project I.", "In Alliance Y 1 Firm chose Project II.", "In Alliance X 1 Firm chose Project I.", "In Alliance X 1 Firm chose Project II.", and "Your final profit for this round is: 250". At the bottom right, there is an "OK" button.

When ALL subjects have clicked OK (alternatively, when 40sec have passed) subjects are brought back to the decision task and period 2 starts with the following message.



Profit calculator: the profit calculator was shown on a screen on another monitor to the right of the monitor showing the decision task.

Profit Calculator

Click on the option buttons to make any hypothetical investment choice between Project I and Project II for any firm at any time.

The project that is hypothetically chosen by each firm is shown in the green box corresponding to that firm.

The profit that each firm would accrue from such a combination of choices is shown in the yellow box corresponding to that firm.

Please note that this is only a calculator - the "choices" and "profits" here are strictly hypothetical.

In Alliance X

Firm A chooses

Project I ☐ Project II ☐

Project I

In Alliance Y

Firm B chooses

Project I ☐ Project II ☐

Project II

Firm C chooses

Project I ☐ Project II ☐

Project I

Firm D chooses

Project I ☒ Project II ☐

Project I

and

and

and

Firm A earns

250 points

Firm B earns

100 points

Firm C earns

125 points

Firm D earns

125 points

and

and

and

32

A3. Further dynamic analysis

We denote rivals' decisions as x_C and x_D . All our theories are based on $x = (x_A, x_B, x_C + x_D)$, apart from *PD* where (x_A, x_B) . Taking all four games together, this leads to forty different decision situations with forty frequencies of contributions in the next round. The number of degrees of freedom for chi-square goodness of fit tests is forty minus the number of estimated parameters. Table A1 shows the empirical distribution of contributions in the forty decision situations.

Table A1. Frequency distribution of decisions *Coop* = 0 or 1 dependent on the vector of previous decisions $x = (x_A, x_B, x_C + x_D)$

Game	x_B	$x_C + x_D$	$x_A = 1$		$x_A = 0$	
			<i>Coop</i> = 1	<i>Coop</i> = 0	<i>Coop</i> = 1	<i>Coop</i> = 0
<i>PD</i>	0		9	26	16	76
<i>PD</i>	1		43	11	7	28
<i>INFO</i>	0	0	2	8	3	17
<i>INFO</i>	0	1	2	4	6	14
<i>INFO</i>	0	2	5	9	11	21
<i>INFO</i>	1	0	31	1	0	10
<i>INFO</i>	1	1	24	4	2	4
<i>INFO</i>	1	2	19	5	2	12
<i>RIVALRY</i>	0	0	2	15	10	94
<i>RIVALRY</i>	0	1	4	8	7	27
<i>RIVALRY</i>	0	2	0	2	1	5
<i>RIVALRY</i>	1	0	3	3	4	13
<i>RIVALRY</i>	1	1	2	2	4	8
<i>RIVALRY</i>	1	2	0	0	0	2
<i>SYNERGY</i>	0	0	4	17	9	27
<i>SYNERGY</i>	0	1	11	19	7	35
<i>SYNERGY</i>	0	2	2	3	2	4
<i>SYNERGY</i>	1	0	3	3	5	16
<i>SYNERGY</i>	1	1	7	3	11	19
<i>SYNERGY</i>	1	2	2	2	2	3

Behavioral theories used in the repeated prisoner's dilemma

Dal Bó and Fréchette (2018) survey the performance of simple behavioral rules as Tit-for-Tat (*TfT*), Grim, Win-Stay-Lose-Shift (*WSLS*), “never contribute” (0), “always contribute” (1) and some other less often played rules. We extend these five rules by *Inertia* = “repeat your previous choice” and *TfToth* = “contribute if $x_3 + x_4 = 2$ “or, alternatively, $x_3 + x_4 \geq 1$. In addition, we assume a “trembling hand” probability ε which describes random deviations from the rules (same for all rules). With these components we estimate population shares (*PopSh*, the share of *TfToth* is the residual share) for the rules. The result is presented in Table A2.

Table A2. Estimation of a mixture of behavioral rules. *PopSh* = population share*

	<i>l</i>	<i>0</i>	<i>TfT</i>	<i>Grim</i>	<i>WSLS</i>	<i>Inertia</i>	<i>TfToth</i>	ε
<i>PopSh.</i> , ε	0	0	0.134	0.535	0.067	0.265	0	0.160
<i>S.D.</i>	0.005	0.015	0.107	0.464	0.067	0.016	-	0.030

* $LL = -473.6$. Standard deviations in this and all following tables are computed from the Hessian of the loglikelihood function (if non-singular).

Table A3. Estimated parameters and standard deviations for altruism and FS models with exogenous beliefs and constant coefficients of utility functions*

	λ	β_B	α_B	β_{out}	α_{out}	η	δ	AIC
<i>Alt</i>	3.761	0.358		1		0	-	1059
<i>S.D.</i> *	-	-		-		-	-	
<i>Alt-LLA</i>	6.930	-0.075		0.601		0.041	0.578	892.8
<i>S.D.</i>	1.435	0.204		0.170		0.089	0.072	
<i>RDAhist</i>	2.607	1.000	0.461	0.968	0.995	0.208	-	923.8
<i>S.D.</i>	0.402	0.015	0.066	0.200	0.223	0.092	-	
<i>RDAexp</i>	3.735	1.000	0.602	0.697	0.724	0.055	-	900.4
<i>S.D.</i>	0.504	0.010	0.035	0.141	0.145	0.086	-	
<i>RDAhist-LLA</i>	5.728	0.152	-0.005	0.705	0.732	0.099	0.543	893.4
<i>S.D.</i>	1.665	0.209	0.210	0.223	0.237	0.111	0.083	
<i>RDAexp-LLA</i>	9.69	-0.434	-0.217	0.417	0.422	0.017	0.551	896.2
<i>S.D.</i>	4.375	0.501	0.264	0.222	0.226	0.099	0.050	
<i>FS</i>	1.591	1.000	1.000	1.000	1.108	0.101	-	939.8
<i>S.D.</i>	0.164	0.007	0.011	0.011	0.718	0.118	-	
<i>FS-LLA</i>	6.715	-9.631	0.153	0.168	0.167	0.008	0.504	918.0
<i>S.D.</i>	1.169	34.33	0.121	0.183	0.181	0.084	0.020	

* $\alpha_{out} = \alpha_C = \alpha_D$ and $\beta_{out} = \beta_C = \beta_D$.

Table A4. Estimated parameters and standard deviations for reciprocity models*

	λ	β_B	α_B	θ_B	β_{out}	α_{out}	θ_{out}	δ	AIC
<i>DynAlt</i>	9.980	-0.182		-3.504	0.360		-0.137	-	517.7
<i>S.D.*</i>	-	-		-	-		-	-	
<i>DynFS=CR</i>	17.21	0.506	0.287	0.169	2.526	0.894	0.000	-	467.6
<i>S.D.</i>	9.051	0.081	0.073	0.189	0.337	0.111	0.005	-	
<i>DynAlt-LLA</i>	1.575	0.069		1.696	1.096		0.525	0.443	480.6
<i>S.D.</i>	0.228	5.486		132.4	0.377		0.287	0.111	
<i>CR-LLA</i>	17.24	-0.507	0.287	0.168	2.527	0.894	0.000	0.000	467.6
<i>S.D.</i>	1.307	0.027	0.020	0.223	0.144	0.039	0.005	0.009	

* *DynAlt* with changing altruism coefficients and *DynFS = CR* with changing *FS* coefficients

with endogenous beliefs (QRE). $\alpha_{out} = \alpha_C = \alpha_D$, $\beta_{out} = \beta_C = \beta_D$, and $\theta_{out} = \theta_C = \theta_D$.

A4. Sociocognitive effects on competitive behavior

Regulatory focus effects on conflict

Competitive effects can also arise from concerns for individual gain or loss. We therefore also explore the sociocognitive perspective of Regulatory Focus Theory (RFT; Higgins, 1998). RFT asserts that individuals are motivated along two distinct dimensions, and that value is derived not just from reaching specific goals but also from the means by which they are reached. Promotion focus is concerned with the pleasure of gain and pain in non-gain, and relates to accomplishments, hopes, and aspirations. It pursues “gain attainment goals” via “approach strategic means.” Prevention focus is concerned with the pleasure of non-loss and pain of loss, and relates to safety, responsibility, and obligations. It pursues “loss avoidance goals” via “avoidance strategic means.” Promotion and prevention focus map onto competitive gains and losses, respectively, in intergroup conflict.

We postulate that cooperation increases (decreases) with promotion (prevention) focus, because of the welfare gain (loss) from partner cooperation (defection), which outweighs the prospect of non-gain (non-loss) from partner defection (cooperation). Also, the positive effect

of promotion (prevention) focus on intragroup cooperation should be more (less) pronounced when we increase the gains from cooperation, i.e. synergy. In contrast, prevention focus (promotion) should be more (less) sensitive to potential losses from intergroup conflict when rivalry increases, and in turn we expect intragroup cooperation to increase, especially when monetary losses from rivalrous effort are realized. The same argument holds for psychological losses from observed outgroup effort, even without payoff-consequences.

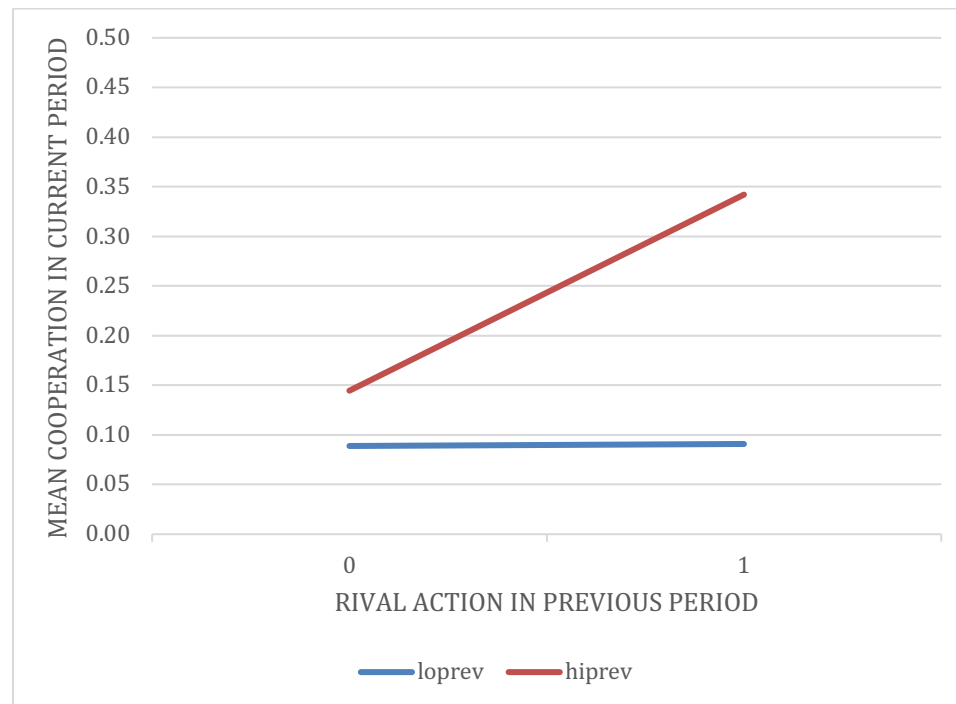
Testing regulatory focus effects

We measure individual RF along two psychometrically distinct and independent subscales using the regulatory focus questionnaire (RFQ; Higgins et al., 2001). The RFQ bases its measures on an individual's self-reported subjective history of success in attaining promotion and prevention focus goals. They show that subjects with subjective histories of success with promotion (prevention) focus are more inclined to promotion (prevention) focus. Subjects had to rate the frequency of having encountered 11 different types of events on a Likert scale between 1 ("never or seldom" or "never true") and 5 ("very often" or "very often true"). The RFQ measures chronic regulatory focus based on prior life experiences, for example, "I have found very few hobbies or activities in my life that capture my interest or motivate me to put effort into them." "I feel like I have made progress toward being successful in my life." "How often did you obey rules and regulations that were established by your parents?"

RF raw scores are computed following Higgins et al. (2001).⁸ Figure A1 shows higher levels of cooperation in *RIVALRY* for High-Prevention compared to Low-Prevention individuals (classified by median split), indicating a higher degree of pre-emption. It also shows cooperation as a function of rival action in the previous period, and Low-Prevention individuals do not demonstrate responsiveness to observed rivalry the High-Prevention individuals do.

⁸ The key for RFQ is as follows: Promotion = (6 - rfe1) + rfe3 + rfe7 + (6 - rfe9) + rfe10 + (6 - rfe11), and Prevention = (6 - rfe2) + (6 - rfe4) + rfe5 + (6 - rfe6) + (6 - rfe8), then mean centered them to avoid undesired multicollinearity (Marquardt, 1980). Data for four subjects were dropped in the regressions due to human error resulting in the loss of their RFQ data.

Figure A1. Cooperation as a function of prevention focus and rival action in the previous period in *RIVALRY**



* There were only 5 observations of both rivals cooperating in the previous period and so these estimates are not presented graphically. Based on the median split we have two groups: one with low prevention focus (loprev) and another with high prevention focus (hiprev).

Model 3 (Table A4) shows that promotion focus has a positive albeit marginally significant effect on increasing cooperation in response to rivalrous effort ($LRivalCoop \times Promotion$). This suggests that partners are more inclined to collaborate to achieve gains that abate non-gains from rivalry. Model 4 shows that prevention focus increases cooperation in response to anticipated competitive threats as rivalry increases ($r_{out} \times Prevention$), and to observed rivalrous effort $LRivalCoop$ ($LRivalCoop \times Prevention$). Consistently, it makes the decision maker less prone to the attenuation effect of increased synergy r_{in} ($r_{in} \times Prevention$).

Table A4. Binary logit regressions of cooperation*

	(3)	(4)
<i>Round</i>	-0.135*** (0.032)	-0.156*** (0.040)
<i>LPartnerCoop</i>	1.129*** (0.280)	0.905*** (0.231)
<i>r_{in}</i>	2.820* (1.678)	2.974 (2.114)
<i>r_{out}</i>	-7.541** (3.056)	-4.943 (3.413)
<i>LRivalCoop</i>	0.108 (0.235)	0.190 (0.212)
<i>Promotion</i>	0.410 (0.395)	
<i>r_{in} × Promotion</i>	-0.629 (0.579)	
<i>r_{out} × Promotion</i>	0.189 (0.636)	
<i>LPartnerCoop × Promotion</i>	0.001 (0.077)	
<i>LRivalCoop × Promotion</i>	0.121* (0.063)	
<i>Prevention</i>		0.382 (0.300)
<i>r_{in} × Prevention</i>		-1.181** (0.468)
<i>r_{out} × Prevention</i>		2.633*** (0.884)
<i>LPartnerCoop × Prevention</i>		-0.112 (0.077)
<i>LRivalCoop × Prevention</i>		0.148** (0.060)
<i>Constant</i>	-1.700 (1.228)	-2.159 (1.450)
Pseudo <i>R</i> ²	0.1471	0.1864
<i>N</i>	612	612

* Robust standard errors in parentheses, clustered by session. The p-values reported here are two tailed, with * denoting $p < 0.10$, ** denoting $p < 0.05$, and *** denoting $p < 0.01$.

Discussion of regulatory focus effects on conflict

This negative effect of conflict on cooperation, however, is dampened by loss avoidance as captured by prevention focus, which induces competitive effort when provoked by rivals: cooperation with partners increases in response to both anticipated and realized competitive threats. To frame such behavior in RFT terms, assume the NE payoff as a benchmark, being the *ex ante* expected payoff. From a promotion (prevention) perspective, an outcome is a gain (loss) if the payoff is more (less) than the NE payoff – otherwise the outcome

is a non-gain (non-loss). Thus, relative to the equilibrium payoff from mutual defection, mutual cooperation yields mutual gains while unilateral cooperation (defection) yields individual losses (gains).

Prevention focus secures against losses and ensures non-losses. Intragroup cooperation might trigger off an intergroup war if rivals retaliate in kind, but defection cannot “insure” against losses inflicted by rivals. To prevention focus, there is no “safety” in defection especially if rivals have acted aggressively and further attacks are expected. Further attacks, if not pre-empted, cause loss and in turn regret (Zeelenberg et al., 2002). In contrast, cooperation can help alleviate losses from rival effort, i.e. net non-loss. Retaliation is in line with the sense of duty and pre-emptive tendencies (Freitas et al., 2002) of prevention focus, which motivates vigilance in the face of threats and in response to realized threats, i.e. “true negatives” (Crowe and Higgins, 1997).

The intensity of inflicted “loss” under prevention focus is higher than that of an equivalent “non-gain” under promotion focus and this motivates retaliation. Thus, while conflict should still decrease the likelihood of intragroup cooperation (so as to avoid intergroup conflict), this effect is weakened by prevention focus especially when threats are realized. Finally, the increased expected payoff from greater synergy is perceived as a “non-loss,” which is less salient than the risk of “losses” from partner defection and triggering off a war with rivals. Thus, its increasing effect on intragroup cooperation is weakened by prevention focus.

To conclude, with rivalry, ingroup cooperation increases (decreases) with prevention focus (prosocials). When anticipated or provoked, prevention focus induces cooperation to fend against group conflict. Overall, groups are entities that passively resist competitive threats.