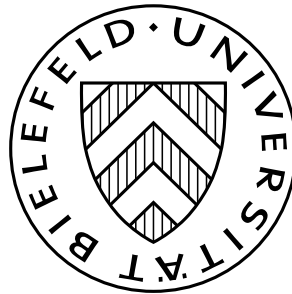


March 2011

A Dynamic Model of Reciprocity with Asymmetric Equilibrium Payoffs

Niko Noeske



A Dynamic Model of Reciprocity with Asymmetric Equilibrium Payoffs

Niko Noeske*

Abstract

We analyze indirect evolutionary two-player games to identify the dynamic emergence of (strong) reciprocity in a large number of economic settings. The underlying evolutionary environment allows for an arbitrary initial population state provided that every degree of the compact space of reciprocity is adherent to at least one individual of the corresponding continuum population. The basic results, which essentially maintain the evolutionary viability of reciprocity, are, in several directions, context dependent, and minimum valid for the wide class of evolutionary dynamics which hold for regularity and payoff-monotonicity. The evolutionary solution concept which is applied to elevate the explanatory power of emerging Nash equilibria is dominance solvability, in this case, for continuous strategy spaces. An asymmetric aspect comes into play since the actions of the evolutionary players are not only determined by the current state of reciprocity but also by their inherent, context-free preferences towards others which differ among one another devoid of being endogenized in the time span of the dynamic process at hand.

Keywords: Reciprocity; Evolutionary Game Theory; Dominance Solvability; Asymmetric Game Setting; Payoff-monotonic Dynamics

*Institute of Mathematical Economics, Bielefeld University.

1 Introduction

In recent years, it has been established that the orthodox assumption on material or monetary¹ selfishness of players is not necessarily sustainable in economic modeling. While the assumption of exogenous given selfishness fits fairly well in some economic contexts,² real life evidence and experimental data suggest that people do not behave consistent with this postulate in general. Most convincing studies include ultimatum/dictator games (e.g. Güth et al., 1982; Andreoni and Miller, 2002; Camerer, 2003), or public goods contribution games (e.g. Andreoni et al., 2002).

The present research contributes to the literature by analyzing a dynamic model of reciprocity in an evolutionary framework.³ In essence, reciprocity refers to peoples' desire to reward perceived kindness and to punish perceived unkindness. The present form of reciprocity invokes the idea that subjective values include to be sensitive to opponents' intrinsic preferences. More precisely, our suggestion of players' *psychological payoff* including other-regarding motivations follows the model of Levine (1998) who demonstrates a striking consistency with results from experimental lab studies. Levine shows that his model is useful to understand several results from ultimatum games and market experiments. Formally, player i seeks to maximize her subjective well-being, given by⁴

$$v_i = u_i + \sum_{j \neq i} \frac{\alpha_i + \lambda \cdot \alpha_j}{1 + \lambda} \cdot u_j,$$

where u_i and u_j are material payoffs, $\alpha_i, \alpha_j \in (-1, 1)$ are social preference parameters, and $\lambda \in [0, 1]$ symbolizes a weight which player i puts on player j 's preference. Accordingly, individuals are not only concerned with their own material payoff but also with that of their opponents. This fact is in the first place due to intrinsic preferences like altruism. However, by considering an extra dimension of reciprocity⁵, the individual weight which is

¹Putting the meaning of selfishness to an economic environment is essential for the present study and related models with so-called *other-regarding* preferences. To make this point, consider Joel Sobel's comment on the hypothesis *only the selfish survive*: "with sufficient freedom to define "selfish" this statement is a tautology" (Sobel (2005, p. 430)). In other words, motivations may depend on others' motivations but their exclusive personality is a banality in the final analysis. This thinking appears trivial yet corresponds to a famous theory called "psychological egoism" which claims that anything we do for others is just because of increasing our own welfare.

²Generally, in highly competitive settings with many players and with one-shot and/or anonymous play (e.g. financial markets), one can assume that the average behavior reflects a high degree of material greed.

³As described in the next section, one should be aware of the fact that reciprocity has different definitions in the economic literature. See also the survey paper by Sobel (2005) for this issue.

⁴This functional is the exact writing of Levine (1998, p. 597). Sethi and Somanthan (2001) introduce a similar model by replacing player i 's weight on player j 's material profit $[\alpha_i + \lambda \cdot \alpha_j] / [1 + \lambda]$ by $[\alpha_i + \lambda \cdot (\alpha_j - \alpha_i)] / [1 + \lambda]$, where $0 \leq \alpha_i < 1$ and $\lambda \geq 0$. Their specification allows an altruist to place a negative weight on a selfish individual which is not possible under Levine's definition. Apart from the denotative conception of reciprocity and for the sake of technical simplicity, the precise definition of the players' subjective payoffs are once more slightly modified in the present study.

⁵Although Levine does not explicitly connect with the term "reciprocity" in his study (however, Sethi and Somanthan do), the parameter λ clearly represents much of the features of reciprocity in any environment of non-anonymous interaction.

placed on the opponents' profit varies additionally with respect to the opponents' intrinsic preference. The advantage of the well-being functional with *two* preference dimensions is that it allows us to explore why people sometimes behave contrary to their true attitudes or *ethos*. In particular, the daily observance, whether in economic or other social life settings, of intrinsic good people behaving badly (or selfishly) sometimes or intrinsic selfish (or bad) people behaving well sometimes centers the motivation of the current study.

Technically, we use an indirect evolutionary environment to identify the cultural viability of reciprocity in a broad class of pre-programmed populations where the players engage in many different types of strategic two-player games. Indirect evolution allows the players to choose a behavior which follows their perceptive payoffs but the players receive evolutionary fitness (or reproductive success) according to their “true”, objective payoffs. The presupposition that players aim to maximize their own idiosyncratic preferences but that economic success “regulates the market” is by now a central tenet in economic modeling and successfully opposes the neo-classical feature of pure economic selfishness. Consequently, a reciprocity disposition which yields a larger objective payoff tends to become more prevalent in a certain population while dispositions with relative low objective payoffs tend to decrease. The idea that social preferences beside selfishness are profitable in some strategic environments is well documented in the relevant literature and dates back at least as far as Schelling (1960). The basic reason for this is that (at least somewhat) recognized preferences provide a commitment device which makes unselfish players potentially the better performers in interdependent settings. Another precursor of this theme is Frank (1987, 1988) who finds that recognizable emotions have the strategic power to change the actions of others, and therefore features the ability to increase the profit of the possessors. Güth and Yaari (1992) then formally introduce the approach of indirect evolution. In fact, they use this approach to challenge a form of reciprocity by having regard for individuals who have tendencies for rejecting unfair offers in ultimatum bargaining. However, the reciprocity motive is somewhat restricted there because the agents are not able to feel subjective benefits from proposing fair distributions. In line with related literature, the finding of Güth and Yaari is that the observability of types guarantees that reciprocators gain from their attitudes in material terms while opportunists relatively loose.⁶

In the present study, we analyze the evolving of reciprocity in the wide frame of regular and payoff-monotonic selection processes. For that purpose, we use a result which is shown by Heifetz et al. (2007): if the evolutionary game on the level of biases (which is located at equilibrium behavior which results from the players' idiosyncratic well-being functionals) is dominance solvable, in the sense of Moulin (1984), for continuous preference spaces, then, the limiting population can be characterized under any payoff-monotonic selection dynamics. The replicator dynamics for general distributions (cf. Oechssler and Riedel (2001) for a detailed technical survey) is subsumed by the class of payoff-monotonic growth-rate

⁶Huck and Oechssler (1999) illustrate viability of preferences for rejecting unfair divisions even if preferences are unobservable provided that the population is small and that the distribution of preferences is known in the population.

functions that is applied to the present study. The more general class considered here explicitly allows to interpret the evolving of reciprocity not just on a biological level but rather on a cultural level of learning (including imitation) or education where a reciprocity disposition that is adequate to achieve higher material returns is replicated faster. Note that our approaching is somewhat enhanced in comparison with the practice of a group of research articles that assumes the evolving of other-regarding issues (like reputation, social preferences, positional goods, ideologies, and so on) on a biological level and treats the results simply as a metaphor to interpret the dynamics in terms of cultural spreading.⁷ Yet, this is not the only reason why the present methodological approach is prepared for exploring reciprocity. In alignment with the dynamical system, we are allowed to assume *any* arbitrary initial population distribution provided that it is described by a compact interval of reciprocity. This technical feature equips the results with a striking general character.

A first observation of our work shows that a sufficient reciprocal player who is intrinsically spiteful (altruistic) may behave benevolent (spiteful) if the opponent player is sufficient contrary in the intrinsic attitude and the strategic situation requires that kind of behavior. In fact, the sign of the players' overall concern for the other players' payoff is always determined by the strategic situation, i.e. the players show a positive concern for the other players if strategic complements are present and a negative concern if strategies are substitutes—a result which is reminiscent to related work of [Possajennikov \(2000\)](#) and [Bester and Güth \(1998\)](#). A further observation says that if player A's intrinsic preference lies below a certain threshold (if player A is relatively spiteful), then player A's tendency to reciprocity is higher the lower the intrinsic preference of player B (the nastier player B)—if player A's intrinsic preference lies above the threshold, then player A's tendency to reciprocity is higher the higher player B's intrinsic preference is. The threshold depends on the strategic type of the underlying game.

The remainder of the paper works as follows. The first part of section 2 surveys some recent approaches of reciprocity or fairness in economics which are, in some obvious sense, connected with the present study. The second part illustrates why our treatment of subjective payoff perception, involving reciprocity, has the strategic power to resolve social dilemma conflicts. Section 3 then introduces the basic model under which the viability of reciprocity is analyzed. Section 4 concludes. Technical details about the applied dynamics, and a figure and a table, that specify some initial conditions, appear in the appendices.

⁷Answering the question whether preference evolution relies on a biological or cultural selection process is a subtle task and rarely elucidated in much detail in related work. However, in our case, it is indeed useful to examine reciprocity on a rather short-termed cultural level since we assume two differently treated preference dimensions which initiate the equilibrium actions: one (altruism/spite) is exogeneous given and not evolving while the other (reciprocity) is endogenized and evolving. With these specifications, it is appropriate to think of reciprocity as a cultural norm which changes within the time span of cultural spreading while altruism/spite changes more seldom by gene transmission. For a deeper understanding of the different selection processes (and the associated speed differences) in evolutionary game theory, see [Selten \(1991\)](#).

2 On Reciprocity

2.1 Related Models and Current Treatment

In order to explain the emergence of other-regarding preferences, there are by now several studies that try to identify more or less complicated models which give insight into the economic psychology of people in strategic situations. The common theme in these models is the antithesis to the neo-classical assumption that people's behavior is thoroughly driven by material selfishness. In order to narrow the wide spectrum of recent approaches and to connect with the present work, it is useful to concentrate on prominent models which explicitly incorporate a certain motive of reciprocity or *fairness*.⁸ One such class assumes that subjective benefits are motivated from inequity aversions of own and other's economic gains. Put differently, the players' actions are initiated by distribution considerations. [Fehr and Schmidt \(1999\)](#) propose for this approach. In their regard, the subjective well-being functional of player i in the standard two-player setting is formalized as

$$U_i = \pi_i - \alpha_i \max[\pi_j - \pi_i, 0] - \beta_i \max[\pi_i - \pi_j, 0],$$

where π_i, π_j are material payoffs and $\alpha_i \geq \beta_i \geq 0, \beta_i < 1$ are weight parameters. Hence, the perceptive payoff of player i differs significantly in the issue whether she and her opponent j are approximately equal rich in economic values; viz., the players are pre-programmed to feel satisfaction from being about as rich as the opponent players. To further summarize, the players are inequity-averse ($\alpha_i, \beta_i \geq 0$), dislike inequity more if it springs from own relative loss ($\alpha_i \geq \beta_i$), but like gaining profit more than reducing inequity ($\beta_i < 1$). A similar model, also motivated by the idea that the players aim to reduce inequity in their material payoffs, is developed by [Bolton and Ockenfels \(2000\)](#). In the two-player setting, Bolton and Ockenfels assume that the personal well-being of player i is determined via the (possibly non-linear) term

$$U_i = v_i \left(\pi_i, \frac{\pi_i}{\pi_i + \pi_j} \right),$$

where $v_i(\cdot, \cdot)$ is globally non-decreasing, concave in the first argument (the material payoff of player i), and strictly concave in the second argument (the relative material payoff of player i). The models of Fehr/Schmidt and Bolton/Ockenfels are both motivated by the idea that players act according to satisfy their fairness emotions by reducing economic inequity. However, the players in these models are distributional motivated and do not explicitly estimate the individual types of the opponents, i.e. the players do not differentiate in the other players' intentions or preferences. Inspired by the psychological

⁸Usually, the will to reciprocate springs from the will to being fair. However, the concept of fairness is likewise of somewhat ambiguous use in economics. At this point, we refrain from a broader discussion on fairness and point to the well-being functionals of this section for examining the specific conception.

game-theoretical approach of [Geanakoplos et al. \(1989\)](#), there is a somewhat more complex class of fairness models which accounts for these elements and seems to reflect reality more detailed. In psychological games, the players' preferences depend on their beliefs about the other players' intentions. By using normal form games, [Rabin \(1993\)](#) proposes for this technique. With his notation, he assumes that individual i plays according to her expected utility

$$U_i(a_i, b_j, c_i) = \pi_i(a_i, b_j) + \tilde{f}_j(b_j, c_i)[1 + f_i(a_i, b_j)],$$

where a_i, b_j , and c_i are, in this order, the strategy chosen by player i , the belief of player i about the strategy of player j , and the belief of player i about the belief of player j about the strategy of player i . $\pi_i(a_i, b_j)$ is player i 's material payoff, $\tilde{f}_j(b_j, c_i)$ is player i 's belief about the kindness of the opponent j towards player i , and $f_i(a_i, b_j)$ symbolizes the kindness of player i towards player j . If equilibrium play is reached, the players' beliefs about the other players' intentions are true and the players base their actions on these beliefs and the subsequent actions of the other players. In line with Rabin's approach but with the purpose to expand to extensive form games, [Dufwenberg and Kirchsteiger \(2004\)](#) assumes a similar model. The basic difficulty of the extensive form is given by the fact that the players have to adjust their perceived utility by updating their beliefs about the others' intentions at each node of the sequential game tree to ensure useful results. Contrary, in [Geanakoplos et al. \(1989\)](#) and [Rabin \(1993\)](#), the players have only initial beliefs about the others' intentions which make the equilibrium analysis easier. [Falk and Fischbacher \(2006\)](#) and [Charness and Rabin \(2002\)](#) propose equilibrium models which basically incorporate both aspects the distributional one and the intention-based. However, opposing to the approach used in the current paper, the applicability of these models is somewhat limited which is basically due to the assumption of higher order beliefs about the others' intentions and the appearance of many equilibria. Levine's model and the version used here depict a third way of modeling reciprocity. In particular, the subjective utility functions of the players depend on the beliefs about the intrinsic preferences of others, i.e. the players reciprocate to the perceived preference of the respective opponent player.

From the previous sentences, it becomes clear that the meaning of reciprocity is ambiguous in terms of functional forms subsuming a motive of fairness. However, discriminations of reciprocity are multi-dimensional existent. In the following, we will mention some further basic facets which appear repeatedly throughout the literature. One prominent aspect is to distinguish between weak and strong reciprocity. Weak reciprocity stands for the conception that people reciprocate in order to gain higher material returns in the future by sustaining collaboration. Typically, weak reciprocity relies on reputation and repeated interaction in orthodox economic modeling and is basically not different from pure selfishness in social preference terminology. In a key paper, [Trivers \(1971\)](#) uses the term *reciprocal altruism* which is identical to weak reciprocity for which he shows sustain-

ability under infinite repeated interactions.⁹ In contrast, strong reciprocity refers to the conception that people show cooperative or retaliatory behavior, even if there is no reason to expect higher material returns in the future (e.g. Gintis, 2000). Moreover, people are willing to sacrifice own profit in order to either help friends or harm enemies. Under this aspect, strong reciprocity is *really* other-regardingly intended.¹⁰

The differentiation of positive and negative reciprocity among the literature is self-explanatory (e.g. Hoffmann et al., 1998). Loosely speaking, positive reciprocity describes the tendency to reward kind people while negative reciprocity describes the tendency to harm cruel people.

Another common aspect is to differentiate between direct and indirect reciprocity (e.g. Nowak and Sigmund, 2005). Direct reciprocity describes the routine that if ‘person A helps (harms) person B, then person B helps (harms) person A’ while indirect reciprocity states that if ‘person A helps (harms) person B, then person C helps (harms) person A’.

The specific notion of reciprocity used in the current paper is determined by the subjective well-being functionals of section 3 and the underlying methodological approach, and, in this regard, described as follows.

In comparison with the significant recurrent features of economic reciprocity, the present shape exhibits the following attributes:

- *preference-based*
- *strong (intrinsic)*
- *both, positive or negative*
- *indirect.*

Though, the stated attributes are not intended to identify an objective, “true” definition of reciprocity in economics. More precisely: the aim of this paper is not to elucidate the meaning of reciprocity how it should be used in economics but rather to assume a form of reciprocity that exhibits some predominant features which appear frequently in the literature, and to explore under what circumstances this form of reciprocity can survive in an evolutionary process.

⁹As already discussed in the literature, the denotation “reciprocal altruism” appears somewhat inadequate in this respect, cf. Hoffmann et al. (1998, p. 338). They argue convincingly that “I am not altruistic if my action is based on my expectation of your reciprocation”. Another thread of research comments that weak reciprocity is not reciprocity and would therefore probably be unhappy with Trivers’ denotation even in this aspect. For example, Fehr and Fischbacher (2002, C3) write: “It is important to emphasize that reciprocity is not driven by the expectation of future material benefit. It is, therefore, fundamentally different from “cooperative” or “retaliatory” behaviour in repeated interactions.”

¹⁰For obvious reasons, Sobel (2005) substitutes “strong” with “intrinsic”, and “weak” with “instrumental”.

2.2 A Simple Illustration

As we will see, the implications of strong reciprocity which we obtain in our basic model are somewhat intricate to retrace (however, the trend and interpretation of these results remain on a plain level). So, the following formulation of the symmetric two-player prisoners' dilemma is intended to give a simple illustration why the current treatment of subjective payoff perception, involving reciprocity, has the strategic power to resolve social dilemma conflicts and overcome spite.¹¹ Consider the following 2×2 matrix.

	cooperate	defect
cooperate	$\xi - v, \xi - v$	$-v, \xi$
defect	$\xi, -v$	$0, 0$

Figure 1: A Prisoners' Dilemma

As is the rule in the matrix design, one player is the row player and the other plays the column; the first number in each matrix entry is the payoff received by the row player and the second one belongs to the column player. The strategy *cooperate* is connected with a private loss of $v > 0$ and a benefit to the other player of $\xi > v$. The strategy *defect* yields neither a loss nor a benefit. If two intelligent and self-interested individuals play this game exactly once, we face the well-known dilemma where both players *defect* in order to reach a higher payoff regardless of the strategical choice of the other player. However, *cooperate* would be mutually better since both players' outcome is higher under this strategy profile: $\xi - v > 0$ with $\xi > v$.

Under a simple model of natural selection the same dilemma defines the usual situation. If we think of a population with size N consisting of k cooperators, and hence $N - k$ defectors, the reproductive matrix payoffs (or *fitnesses*) are given by

$$f_C = \frac{k-1}{N-1} \xi - v \quad (\text{cooperators})$$

and

$$f_D = \frac{k}{N-1} \xi \quad (\text{defectors}),$$

and the average fitness is determined by $\bar{f} = \frac{k}{N}(\xi - v)$. In any mixed population the defectors reach a higher fitness than the cooperators so that natural selection tends to decline the fraction of cooperators while the fraction of defectors eventually take over the whole population. Hence, without any model arrangement which favors the outcome of cooperation, the dilemma of the one-shot game trivially persists under natural selection where matrix payoffs correspond to fitnesses.

Consider now the notion of player i 's perceived payoff which we use in the following

¹¹The prisoners' dilemma is the leitmotif in Sethi and Somanthan (2003) for surveying economic reciprocity in the evolutionary game theoretic literature. The current version uses a different notion of reciprocity.

chapters, i.e.¹²

$$U_i = \pi_i + \theta_i \cdot \pi_j \text{ with } \theta_i = \alpha_i \cdot \gamma_i + (1 - \alpha_i) \cdot \gamma_j,$$

where $\alpha_i \in [0, 1]$ defines the individual “norm of reciprocity”, and $\gamma_i, \gamma_j \in [-1, 1]$ defines the individual “intrinsic preference”.¹³ Note that if $\gamma_i \neq 0$, then player i puts a non-zero weight on the material payoff of player j (unless the extremely rare case where $\alpha_i \cdot \gamma_i + (1 - \alpha_i) \cdot \gamma_j = 0$ with $\gamma_i \neq 0$; however, even in this case individual i is intrinsically biased but her disposition does not come into effect only because her reciprocity norm and the intrinsic preferences of both players compensate to zero). Hence, we say that player i is biased if $\gamma_i \neq 0$. Further, we say that player i is materialistic if both hold true $\gamma_i = 0$ and $\alpha_i = 1$, since then $\theta_i = 0$, i.e. the material outcome of player i coincides with her perceived payoff. Accordingly, a materialist places no weight on the other players’ payoff while a biased player i places a weight of ρ_i^b on the payoff of a biased player and a payoff of ρ_i^m on the payoff of a materialist, where

$$\rho_i^b = \alpha_i \cdot \gamma_i + (1 - \alpha_i) \cdot \gamma_j; \quad \rho_i^m = \alpha_i \cdot \gamma_i.$$

If two biased player interact, we reach the following payoff matrix.

	cooperate	defect
cooperate	$\xi - v + \rho_i^b (\xi - v), \xi - v + \rho_j^b (\xi - v)$	$-v + \rho_i^b \xi, \xi - \rho_j^b v$
defect	$\xi - \rho_i^b v, -v + \rho_j^b \xi$	$0, 0$

Figure 2: A Prisoners’ Dilemma with Biased Players

Provided that $\xi - v + \rho_\bullet^b (\xi - v) > \xi - \rho_\bullet^b v$ and $-v + \rho_\bullet^b \xi > 0$, and thus $\rho_\bullet^b > \frac{v}{\xi}$, *cooperate* is a dominant strategy for both players (the $\bullet$ stands for either i or j). Note that if ① $\alpha_i \cdot \gamma_i + \gamma_j > \alpha_i \cdot \gamma_j$ and ② $\alpha_j \cdot \gamma_j + \gamma_i > \alpha_j \cdot \gamma_i$ then $\rho_\bullet^b > 0$ and $(\textit{cooperate}, \textit{cooperate})$ is potentially a Nash equilibrium, depending on the ratio of benefit and loss. It is easily comprehended that both inequalities ① and ② hold if both players have altruistic feelings towards others, i.e. $\gamma_i, \gamma_j > 0$. But even in the case of different intrinsic preferences, i.e. $\text{sign}(\gamma_i) \neq \text{sign}(\gamma_j)$, the reciprocity motive is apt to keep *cooperate* as the agreed strategy. For example, if player j is moderate spiteful, say $\gamma_j = -0,6$, and player i is perfectly altruistic, $\gamma_i = 1$, a high reciprocity norm of player j , say $\alpha_j = 0,3$, can reverse player j ’s natural will to defect. Note that the ability to reciprocity gives a strong impetus to the game. Even if a player has a strong intrinsic attitude ($|\gamma_i|$ is close to 1) a perfect tendency to reciprocity ($\alpha_i = 0$) will always overcome the origin will to either cooperate or defect if the other player has a contrary intrinsic attitude ($\text{sign}(\gamma_i) \neq \text{sign}(\gamma_j)$) since then $\text{sign}(\theta_i) = \text{sign}(\gamma_j) \neq \text{sign}(\gamma_i)$. Of course, if both players are spiteful, i.e. $\gamma_i, \gamma_j < 0$, then ① and ② never hold and the only rational strategy is always defect.

¹²Cf. Eqs. (1) and Eqs. (3) in section 3. For the sake of simplicity, we assume that all dispositions are perfectly observable.

¹³In our main model in section 3 we exclude *perfect* intrinsic preferences for technical reasons. This is not necessary for the current illustrative purpose.

In the case that a biased player i meets a materialistic player j , the action of the biased player is determined by her intrinsic preference and independent of her reciprocity motive, since the sign of ρ_i^m (and thus the sign of γ_i) induce whether to cooperate or defect. Naturally, a materialist always defects and is in the advantageous free-riding position if the opponent is a cooperating benevolent player. Clearly, the pros and cons of being biased in the matrix PD game continue in a standard population model where the matrix payoffs correspond to fitnesses.

The prisoners' dilemma essentially illustrates the strategic advantages which can result from the reciprocity motive when two biased players interact where one player is sufficient altruistic and the other is intrinsically malevolent but sufficient reciprocal to overcome this attitude. It also demonstrates that two altruists can always resolve the social dilemma but the reciprocity motive is not apt to overcome the will of an altruist to cooperate if the other player is materialistic in the above sense. Hence, the expected evolutionary advantages which results from the reciprocity motive seem to depend heavily on the initial population distribution regarding the intrinsic preferences of the players.

3 Model

In this section we will introduce our model of strong reciprocity, state our main result under the assumption that the players perfectly recognize each others types, and interpret the results.

Let there be a large population of evolutionary agents. At each instant in time a pair of agents is matched at random to play the game $\Gamma_U = (\{1, 2\}, \{x, y\}, \{U_1, U_2\})$ with the aim to maximize their subjective well-being, determined by

$$U_1 = \pi_1 + \theta_1 \cdot \pi_2 \tag{1a}$$

$$U_2 = \pi_2 + \theta_2 \cdot \pi_1, \tag{1b}$$

where π_1 and π_2 are material or “economic” (and therefore interpersonal comparable) payoffs (e.g. money). In order to incorporate a broad variety of strategic situations in this study, we assume that the material payoffs are defined by

$$\pi_1 = x \cdot (l \cdot y - x) + x \tag{2a}$$

$$\pi_2 = y \cdot (l \cdot x - y) + y, \tag{2b}$$

where $x, y \in [0, \infty)$ describe the actions or efforts of player 1 and player 2, respectively. The parameter $l \in (\underline{l} < 0, \bar{l} > 0)$ determines the characteristic nature of the game by measuring the kind and extent of strategical interdependence; the l is further specified as soon as required. The specification of the economic payoffs is sufficiently general to illustrate success since the economic interpretations are extensive. The simplest example is a production game with either negative or positive externalities which is determined by the sign

of l . The externality $l < 0$ represents, for example, a common pool resource game where the players exploit a resource with efforts x and y , respectively. Accordingly, the higher player A's input the lower player B's payoff. Oppositely, $l > 0$ determines a game where a more aggressive behavior of one player increases the payoff of the other player like in public good contribution settings. Alternatively, one can assume oligopolistic competition where the efforts are either firms' quantity choices (in a Cournot market) or price choices (in a Bertrand market).

The variables θ_1, θ_2 symbolize the subjective overall concern for the respective opponents' profit, and have deeper meanings, as specified by

$$\theta_1 = \alpha \cdot \gamma_1 + (1 - \alpha) \cdot \gamma_2 \quad (3a)$$

$$\theta_2 = \beta \cdot \gamma_2 + (1 - \beta) \cdot \gamma_1, \quad (3b)$$

where the parameters γ_1, γ_2 are intrinsic preferences or attitudes, either altruism or spite¹⁴ (γ_1, γ_2 include also material selfishness at the peak of neutrality). The variables α, β identify the dispositions to reciprocity which belong to player 1 and player 2, respectively. While the dimension of altruism and spite is an exogenous trait, the dimension of reciprocity is endogenized in the model. This distinction allows, for example, a player who is rather altruistic inclined to behave spiteful in a reciprocal manner, however, by keeping the true character. Or, a player who is rather spiteful by nature is able to behave benevolent without changing the true character during the course of selection. From these specifications, one should think about reciprocity as a cultural norm or convention which is changeable by the dynamical pressures, and in this line, provides the distinct flexibility in the players' behavior. It is more for the sake of distinctiveness that we will sometimes refer to the parameters γ_1, γ_2 as intrinsic preferences and to the reciprocity-variables α, β as cultural norms or conventions; because, in the sense of subjective motivations which distort the economic greed of the players and initiate their actions, α and β belong to the class of intrinsic preferences, too. This is rather a question of definition.

In accordance with [Levine \(1998\)](#) and [Sethi and Somanathan \(2001\)](#), we avoid initially the somewhat unnatural economic situations in which a player is less (or equally) concerned about herself than about the opponent. Formally, the subjective overall concern for the other player satisfies $|\theta_1|, |\theta_2| < 1$. To ensure this, we impose the following restrictions on the preference components:

$$\begin{aligned} \gamma_1, \gamma_2 &\in A = [-1 + \epsilon, 1 - \epsilon] \\ \alpha, \beta &\in B = [0, 1], \end{aligned} \quad (4)$$

where ϵ is positive and small. Both assumptions are intuitively plausible. The first one allows the players to exhibit a negative ("spite": $\gamma_1, \gamma_2 < 0$), a neutral ("egoism": $\gamma_1, \gamma_2 = 0$), or a positive ("altruism": $\gamma_1, \gamma_2 > 0$) intrinsic preference towards others. The second assumption is even more intuitive and shows that the players evaluate their overall concern

¹⁴Alternatively, one can assume that envy or malevolence is the opposite preference to altruism.

by a convex combination of the own and the other players' intrinsic preference. Note that α , (β) induce reciprocal actions only if α , $(\beta) \neq 1$ since in the case of α , $(\beta) = 1$ the agents' subjective overall concern is independent of the opponents' intrinsic preference. Accordingly, we have the following notion.

Definition 1. *The agents possess a tendency to reciprocity whenever α , $(\beta) \in B \setminus \{1\}$.*

Evidently, the intuition of these assumptions is in line with the restriction of the subjective overall concern towards others, which is finally fixed by $|\alpha \cdot \gamma_1 + (1 - \alpha) \cdot \gamma_2| = |\theta_1|$, $(|\beta \cdot \gamma_2 + (1 - \beta) \cdot \gamma_1| = |\theta_2|) < 1$. More precisely, with intrinsic traits of altruism and spite and the present distribution of reciprocity, the players are completely identified over the compact space

$$\theta_1, \theta_2 \in \Theta = [-1 + \epsilon, 1 - \epsilon]. \quad (5)$$

The game setup is close to the one of [Harrison and Villena \(2008\)](#) but differs significantly in several aspects. First, Harrison and Villena concentrate on game settings which exhibit negative externalities ($\frac{\partial \pi_1(x,y)}{\partial y} < 0$, $\frac{\partial \pi_2(x,y)}{\partial x} < 0$) and strategic substitutes ($\frac{\partial \pi_1(x,y)}{\partial x \partial y} < 0$, $\frac{\partial \pi_2(x,y)}{\partial x \partial y} < 0$). This means that a higher input of player A lowers both the actual payoff and the marginal payoff of player B. From Eqs. (2) and their first and second order derivatives it is easy to see that the present setting represents negative externalities ($\frac{\partial \pi_1(x,y)}{\partial y} < 0$, $\frac{\partial \pi_2(x,y)}{\partial x} < 0$) and strategic substitutes ($\frac{\partial \pi_1(x,y)}{\partial x \partial y} < 0$, $\frac{\partial \pi_2(x,y)}{\partial x \partial y} < 0$) if $l < 0$, and positive externalities ($\frac{\partial \pi_1(x,y)}{\partial y} > 0$, $\frac{\partial \pi_2(x,y)}{\partial x} > 0$) and strategic complements ($\frac{\partial \pi_1(x,y)}{\partial x \partial y} > 0$, $\frac{\partial \pi_2(x,y)}{\partial x \partial y} > 0$) if $l > 0$.¹⁵ Note that with $l = 0$ there is no strategic interdependence so that economic competition becomes “monopolistic”. Consequently, the present model incorporates a much broader class of strategic games.

A second difference regards the evolutionary analysis. While Harrison and Villena use the ESS concept to illustrate the evolutionary viability of reciprocity, we use dominance solvability as proposed by [Heifetz et al. \(2007\)](#). The lack of ESS is that its predictions are only static. ESS does not explore to what level evolution will lead the evolving trait of a certain population but can only tell whether a somehow reached population state is immune to rare “mutations”.^{16,17} Contrary, dominance solvability is useful to establish dynamic results with respect to many initial population states that evolve according to the broad class of regular, payoff-monotonic dynamics.

According to the indirect evolutionary approach, the players maximize their subjective well-being which leads to a second stage game located at equilibrium behavior (or *on the level of biases*). Let this game be symbolized by $\Gamma_f = (\{1, 2\}, \{\alpha, \beta\}, \{f_1, f_2\})$, where f_1, f_2

¹⁵The terminology to characterize the strategic environment was introduced by [Bulow et al. \(1985\)](#) in order to distinguish games with upward sloping best-response functions from those with downward sloping best-response functions.

¹⁶Formally, a strategy x^* is ESS, if either i) $\pi(x^*, x^*) > \pi(x, x^*)$ or ii) $\pi(x^*, x^*) = \pi(x, x^*)$ and $\pi(x, x) < \pi(x^*, x)$ for all mutations $x \neq x^*$, see [Maynard-Smith and Price \(1973\)](#).

¹⁷Another lack of ESS is the insufficiency for characterizing dynamic stability of certain evolutionary dynamics like replicator or BNN with continuous strategy sets (cf. [Hofbauer et al., 2009](#), and some references therein).

are the reproductive success defining fitness functions that can be identified by substituting equilibrium behavior in the economic payoffs (Eqs. (2)), which formally corresponds to

$$f_1, f_2 = \pi_1(x^*, y^*), \pi_2(x^*, y^*),$$

where x^*, y^* are equilibrium strategies of $\Gamma_U = (\{1, 2\}, \{x, y\}, \{U_1, U_2\})$. Thus, let both players maximize their perceived payoffs, i.e. $x^* \in \operatorname{argmax}_x U_1(x, y^*)$ and $y^* \in \operatorname{argmax}_y U_2(x^*, y)$, which defines their reaction functions: $x = \frac{1}{2}(1 + ly(1 + \theta_1))$ and $y = \frac{1}{2}(1 + lx(1 + \theta_2))$. Equalizing the reaction functions identifies the unique equilibrium profile of the game (x^*, y^*) , where

$$x^* = -\frac{\theta_1 l + l + 2}{l^2 + \theta_2 l^2 - 4 + \theta_1 l^2 + \theta_1 l^2 \theta_2} \quad (6a)$$

$$y^* = -\frac{\theta_2 l + l + 2}{l^2 + \theta_2 l^2 - 4 + \theta_1 l^2 + \theta_1 l^2 \theta_2}. \quad (6b)$$

From the equilibrium profile, we can comprehend the strategic influence of player A's regard for player B's payoff on player B's strategy. The strategic influence is consistent with the psychological idea that individuals condition their actions on the perceived types of others and do not act uniformly with each other. At this point, it becomes clear that the relatedness of the other players' type and the own equilibrium action requires a positive degree of recognition. Note again that we have assumed this ability of the players in the perfect sense.

Plugging the equilibrium actions in the material payoff functions leads to the individual fitnesses which are functions of the biases,

$$f_1(\theta_1(\gamma_1, \gamma_2, \alpha), \theta_2(\gamma_1, \gamma_2, \beta)) = \pi_1(x^*, y^*) = -\frac{(\theta_1 l + l + 2)(-l + \theta_1 l - 2 + \theta_1 l^2 + \theta_1 l^2 \theta_2)}{(l^2 + \theta_2 l^2 - 4 + \theta_1 l^2 + \theta_1 l^2 \theta_2)^2} \quad (7a)$$

$$f_2(\theta_1(\gamma_1, \gamma_2, \alpha), \theta_2(\gamma_1, \gamma_2, \beta)) = \pi_2(x^*, y^*) = -\frac{(\theta_2 l + l + 2)(-l + \theta_2 l - 2 + \theta_2 l^2 + \theta_1 l^2 \theta_2)}{(l^2 + \theta_2 l^2 - 4 + \theta_1 l^2 + \theta_1 l^2 \theta_2)^2}. \quad (7b)$$

Eqs. (7), the *equilibrium payoffs*, are the central functionals which measure the prevalence of the different types in the game (the specifications of the dynamic process—where successful types proliferate at the expense of abortive types—are given in the Appendix A).

In the following, we assume that the intrinsic preference of player 1 is not exactly the same as the intrinsic one of player 2, i.e. $\gamma_1 \neq \gamma_2$. This assumption gives the game Γ_f an asymmetric character and is reasonable in order to adopt reciprocity in the model. In the case of $\gamma_1 = \gamma_2$, it would be sufficient to behave according to the intrinsic preference altruism or spite to fulfill the characteristic of reciprocity. Formally, there are now two different populations but with intrinsic traits of altruism and spite selected from the same pool. Somewhat informal, one can imagine that nature picks γ_1, γ_2 from the equal distributed set A “with two hands at once”. Based on the usual asymmetric setting in

evolutionary games (cf. Selten, 1980; Weibull, 1995, pp. 64), let us imagine an ex ante symmetric game, denote $\Gamma_{\Gamma_f}^\gamma$. In this game any intrinsic preference parameter is assigned to each of the players with the same probability. This assumption corresponds to “nature plays first” by allocating γ_1, γ_2 to player 1 and player 2. Relying on Selten’s work, the pair of reciprocity biases (α, β) , where α is associated with γ_1 (i.e. the biases of player 1 are given with γ_1 and α) and β is associated with γ_2 , would be evolutionarily stable in the sense of ESS in the ex ante symmetric game $\Gamma_{\Gamma_f}^\gamma$ if and only if the vector (α, β) describes a strict Nash equilibrium of the asymmetric game Γ_f . However, as mentioned before, the question of interest regards the conception of dominance solvability, and hence, the question of which type pass the dynamic evolutionary pressures under many starting conditions.

The following lemma is useful to identify a dominance solvable trait (cf. Heifetz et al., 2007; Moulin, 1984, Theorem 4).

Lemma 1. In order to check for dominance solvability of a particular trait it is sufficient to compute that

(i) the fitness function is continuous, twice differentiable and strictly concave in the particular trait of each player;

(ii) the slope of each player’s best-reply function is less than 1 in absolute value;

and to argue that

(iii) the particular trait is selected from a compact interval.

To start with, a substantial argument for condition Lemma 1(iii) to be satisfied here gives the following remark.

Remark. *The issue of compactness or completeness of the bias spaces is rather a philosophical question. At best, one should think about the opportunity to “select” the subjective norm of reciprocity as a hypothetical choice rather than an alternative reflecting from a permanent conscious state of mind. Accordingly, the particular values of α, β , and thus θ_1, θ_2 , as emotional devices come into the conscious minds and initiate the actions only if the strategic situation requires it yet the whole spaces examining players’ potentials are present at any time.*

In order to examine the viability of reciprocity, we will base our analyses on the results given with the Theorem of Appendix A and Lemma 1; however, we have to extend the setting somewhat since we assume the game Γ_f to be asymmetric.

To emphasize the asymmetric character of the game consider now two different bias spaces with elements θ_1, θ_2 since $\gamma_1 \neq \gamma_2$, however symmetrical types are also possible, i.e. $\theta_1 = \theta_2$. So, $\theta_1 \in \Theta_1 = [-1 + \epsilon, 1 - \epsilon]$ and $\theta_2 \in \Theta_2 = [-1 + \epsilon, 1 - \epsilon]$ where $\theta_1 = \theta_2$ only if $\alpha \cdot \gamma_1 + (1 - \alpha) \cdot \gamma_2 = \beta \cdot \gamma_2 + (1 - \beta) \cdot \gamma_1$ with $\gamma_1 \neq \gamma_2$. Since the position of the players’

roles is initially by no means (i.e. the players are either in position 1 or in position 2 with equal probability) it is relatively straightforward to construct an ex ante symmetric game setup. Thus, the profile parameter $\kappa = (\theta_1, \theta_2)$ is selected from the compact support $K = \Theta_1 \times \Theta_2 = [-1 + \epsilon, 1 - \epsilon] \times [-1 + \epsilon, 1 - \epsilon]$ and determines the evolving game parameter of the ex ante symmetric game with the distribution $G_t = (G_t^1, G_t^2)$ where G_t^1 corresponds to θ_1 and G_t^2 to θ_2 ; the t will sometimes be dropped for convenience.¹⁸ Then, the ex ante symmetric game payoff of an individual with type $\kappa = (\theta_1, \theta_2)$ competing with an individual of type $\tilde{\kappa} = (\tilde{\theta}_1, \tilde{\theta}_2)$ is defined by

$$f(\kappa, \tilde{\kappa}) = \frac{f_1(\theta_1, \tilde{\theta}_2) + f_2(\theta_2, \tilde{\theta}_1)}{2}. \quad (8)$$

Accordingly,

$$f(\kappa, G) = \frac{f_1(\theta_1, G^2) + f_2(\theta_2, G^1)}{2} \quad (9)$$

is the ex ante fitness to type $\kappa = (\theta_1, \theta_2)$ under the distribution $G = (G^1, G^2)$. Having this game texture allows us to transfer methodological results from [Heifetz and Segev \(2004\)](#) who use an asymmetric game setting which is close to ours in order to identify “the evolutionary role of toughness in bargaining” which gives the name to their essay. Extending the terminology of domination as in the Appendix A to the asymmetric game setting we have that $\tilde{\kappa} = (\tilde{\theta}_1, \theta_2)$ (or $\hat{\kappa} = (\theta_1, \hat{\theta}_2)$) is dominated by $\kappa = (\theta_1, \theta_2)$ in iteration $n + 1$ if for every $\acute{\kappa} = (\acute{\theta}_1, \acute{\theta}_2) \in U_n$ we have $f_1(\theta_1, \acute{\theta}_2) > f_1(\tilde{\theta}_1, \acute{\theta}_2)$ (or $f_2(\theta_2, \acute{\theta}_1) > f_2(\hat{\theta}_2, \acute{\theta}_1)$).

Accordingly, we reach:

Lemma 2. Dominance solvability of the asymmetric game Γ_f with the two players’ payoffs $f_1(\theta_1, \theta_2)$ and $f_2(\theta_2, \theta_1)$ implies dominance solvability of the ex ante symmetric game $\Gamma_{\Gamma_f}^\gamma$ with payoff $f(\kappa, \tilde{\kappa})$.

Using now the Theorem of Appendix A and Lemma 2, we have:

Lemma 3. If both players’ asymmetric game is dominance solvable to $\theta_1^*(\alpha^*, \gamma_1, \gamma_2)$ and $\theta_2^*(\beta^*, \gamma_1, \gamma_2)$, respectively, then the profile $\kappa^* = (\theta_1^*, \theta_2^*)$ is the unique outcome of the ex ante symmetric game $\Gamma_{\Gamma_f}^\gamma$ under any regular and payoff-monotonic selection dynamics.

We are now able to prove our main result.

Proposition 1. Consider the game described above with $l \in [-1/4, 0) \cup (0, 3/5]$ and an extra requirement as given below the proof of this proposition (the requirement specifies

¹⁸Of course, the basic evolving trait is the norm of reciprocity, α (respective β), however, for the sake of clarity, it is sometimes benefiting to think of the overall concern as the evolving trait (consider θ_1, θ_2 as an initial random weighting of α and β).

the sets of γ_1 and γ_2 that we consider in different situations l). Then, any initial full-support of the distribution of biases $G = (G_0^1, G_0^2)$ will converge in distribution towards a unit mass on the pair of $(\theta_1^* = \alpha^* \gamma_1 + \beta^* \gamma_2, \theta_2^* = \beta^* \gamma_2 + \alpha^* \gamma_1)$ with the pair of reciprocity norms

$$\left(\alpha^* = 1/2 + 1/2 \frac{-4l - 4 + \sqrt{\Lambda(\gamma_1, \gamma_2, l)}}{l(l-2)(\gamma_1 - \gamma_2)}, \beta^* = 1/2 + 1/2 \frac{4l + 4 - \sqrt{\Lambda(\gamma_1, \gamma_2, l)}}{l(l-2)(\gamma_1 - \gamma_2)} \right),$$

with

$$\begin{aligned} \Lambda(\gamma_1, \gamma_2, l) = & l^4 (\gamma_1 + \gamma_2 + 2)^2 - 4l^3 ((\gamma_1 + \gamma_2 + 1)^2 - 1) + \\ & 4l^2 ((\gamma_1 + \gamma_2 + 1)^2 - 1 - 4\gamma_1 - 4\gamma_2) + 16l(\gamma_1 + \gamma_2 + 2) + 16, \end{aligned}$$

under any regular and payoff-monotonic dynamics.

Proof. According to the Theorem of Appendix A and Lemmata 1-3, the procedure of the proof is to find an equilibrium profile of the asymmetric game $\Gamma_f = (\{1, 2\}, \{\alpha, \beta\}, \{f_1, f_2\})$, and then check for sufficient conditions regarding dominance solvability. Thus, calculating first order conditions of $f_1(\theta_1, \theta_2)$ and $f_2(\theta_1, \theta_2)$ (cf. Eqs. (7)), i.e. $\frac{\partial f_1}{\partial \theta_1}(\theta_1, \theta_2) = 0$ and $\frac{\partial f_2}{\partial \theta_2}(\theta_1, \theta_2) = 0$, and solving for the biases yield

$$\theta_1 = -\frac{(\theta_2 l + l + 2\theta_2 + 2)l}{\theta_2 l^2 + l^2 - 2\theta_2 l - 2l - 4} \quad (10a)$$

$$\theta_2 = -\frac{(\theta_1 l + l + 2\theta_1 + 2)l}{\theta_1 l^2 + l^2 - 2\theta_1 l - 2l - 4}, \quad (10b)$$

where it is now reasonable to account for

$$\theta_1 = \alpha \cdot \gamma_1 + (1 - \alpha) \cdot \gamma_2 \quad (11a)$$

$$\theta_2 = \beta \cdot \gamma_2 + (1 - \beta) \cdot \gamma_1, \quad (11b)$$

in order to find equilibria on the level of reciprocity. To this end, we plug Eqs. (11) in Eqs. (10), equalize α and β , and solve for the equilibria.¹⁹ Accordingly, we reach

$$\alpha_{\pm}^* = 1/2 + 1/2 \frac{-4l - 4 \pm \Phi(\gamma_1, \gamma_2, l)}{l(l-2)(\gamma_1 - \gamma_2)}$$

¹⁹Note that, mathematically, it does not matter whether we substitute the overall concern for its components in Eqs. (10) or in the fitness functions (Eqs. (7)).

and

$$\beta_{\pm}^* = 1/2 + 1/2 \frac{4l + 4 \pm \Phi(\gamma_1, \gamma_2, l)}{l(l-2)(\gamma_1 - \gamma_2)},$$

where

$$\Phi(\gamma_1, \gamma_2, l) = \sqrt{\Lambda(\gamma_1, \gamma_2, l)},$$

with

$$\begin{aligned} \Lambda(\gamma_1, \gamma_2, l) = & l^4(\gamma_1 + \gamma_2 + 2)^2 - 4l^3((\gamma_1 + \gamma_2 + 1)^2 - 1) + \\ & 4l^2((\gamma_1 + \gamma_2 + 1)^2 - 1 - 4\gamma_1 - 4\gamma_2) + 16l(\gamma_1 + \gamma_2 + 2) + 16, \end{aligned}$$

where $\Lambda(\gamma_1, \gamma_2, l) > 0$ if $\gamma_1, \gamma_2 \in A = [-1 + \epsilon, 1 - \epsilon]$ and $l \in [-1/4, 1]$. Further, by regarding the restrictions on γ_1, γ_2 and the strategic setting l , and by analyzing the result sets of $\alpha_{\pm}^*$ and $\beta_{\pm}^*$, respectively, we find that α_{-}^* and β_{+}^* are no possible solutions since $\alpha_{-}^* \notin [0, 1]$ and $\beta_{+}^* \notin [0, 1]$. In order to prove that $\alpha_{-}^*, \beta_{+}^* \notin [0, 1]$ we have to consider 8 cases, or accordingly, we have to verify 8 conditions (each condition corresponds to one case), denote (I) to (VIII). The $\Phi(\gamma_1, \gamma_2, l)$ is dropped in the following since it is a positive value and of no account in any of the 8 conditions.

Case 1: Assume $\gamma_1 - \gamma_2 < 0$ and $l \in (0, 1]$. Then, $\alpha_{-}^* < 0$ implies that (I): $-4 < -l(l-2)(\gamma_1 - \gamma_2) + 4l$, which is true since the right hand side of (I) is positive.

Case 2: Assume $\gamma_1 - \gamma_2 > 0$ and $l \in (0, 1]$. Then, $\alpha_{-}^* > 1$ implies that (II): $-4 < l(l-2)(\gamma_1 - \gamma_2) + 4l$, which is true since the right hand side of (II) is positive.

Case 3: Assume $\gamma_1 - \gamma_2 < 0$ and $l \in [-0, 25, 0)$. Then, we analyze the same condition as under case 2, i.e. (II)=(III), but with negative l and $\gamma_1 < \gamma_2$; however, this condition holds even under these constraints.

Case 4: Assume $\gamma_1 - \gamma_2 > 0$ and $l \in [-0, 25, 0)$. Then, we analyze the same condition as under case 1, i.e. (I)=(IV), but with negative l and $\gamma_1 > \gamma_2$; however, this condition holds even under these constraints.

Case 5: Assume $\gamma_1 - \gamma_2 < 0$ and $l \in (0, 1]$. Then, $\beta_{+}^* > 1$ implies that (V): $4 > l(l-2)(\gamma_1 - \gamma_2) - 4l$, which is true since (V) $= (-1) \cdot (I)$.

Case 6: Assume $\gamma_1 - \gamma_2 > 0$ and $l \in (0, 1]$. Then, $\beta_{+}^* < 0$ implies that (VI): $4 > -l(l-2)(\gamma_1 - \gamma_2) - 4l$, which is true since (VI) $= (-1) \cdot (II)$.

Case 7: Assume $\gamma_1 - \gamma_2 < 0$ and $l \in [-0, 25, 0)$. Then, we analyze the same condition as under case 6 (respective 2), i.e. (VII)=(VI) $=(-1) \cdot (II)$, but with negative l and $\gamma_1 < \gamma_2$;

however, this condition holds even under these constraints.

Case 8: Assume $\gamma_1 - \gamma_2 > 0$ and $l \in [-0, 25, 0)$. Then, we analyze the same condition as under case 5 (respective 1), i.e. $(VIII)=(V)=(-1) \cdot (I)$, but with negative l and $\gamma_1 > \gamma_2$; however, this condition holds even under these constraints.

Hence, the unique equilibrium profile to be further analyzed is given by

$$\begin{aligned} \alpha_+^* &= 1/2 + 1/2 \frac{-4l - 4 + \Phi(\gamma_1, \gamma_2, l)}{l(l-2)(\gamma_1 - \gamma_2)} = \alpha^*, \\ \beta_-^* &= 1/2 + 1/2 \frac{4l + 4 - \Phi(\gamma_1, \gamma_2, l)}{l(l-2)(\gamma_1 - \gamma_2)} = \beta^*, \end{aligned} \tag{12}$$

where we have dropped the subscripts “+” and “−” for convenience.

According to Lemma 1(i), the next step is to show that Eqs. (7) fulfill the properties of

- (◀) “continuity”,
- (▶) “twice differentiability”, and
- (▼) “concavity”,

with respect to α (respective β). By substituting θ_1 and θ_2 for its components, it is easy to comprehend that ◀ and ▶ are satisfied. To show ▼, review that we have defined the fixed and evolving dispositions by the restrictions

$$\begin{aligned} \gamma_1, \gamma_2 &\in A = [-1 + \epsilon, 1 - \epsilon] \text{ with } \gamma_1 \neq \gamma_2, \\ \alpha, \beta &\in B = [0, 1], \end{aligned}$$

and the convex combinations, which totally identify the players, by

$$\left. \begin{aligned} \theta_1 &= \alpha \cdot \gamma_1 + (1 - \alpha) \cdot \gamma_2 \\ \theta_2 &= \beta \cdot \gamma_2 + (1 - \beta) \cdot \gamma_1 \end{aligned} \right\} \in \Theta = [-1 + \epsilon, 1 - \epsilon]$$

Now, let us first check for concavity of Eqs. (7) in θ_1 (respective θ_2). Note that although we deal with asymmetric fitness functions, it is sufficient to show that one player’s payoff is concave in the overall concern, i.e. $\frac{\partial^2 f_A}{(\partial \theta_A)^2}(\theta_A, \theta_B) < 0$ (for $A \in \{1, 2\}$ and $B = 3 - A$), since the pools of θ_A, θ_B are equal. Calculating the second derivative with respect to Eqs. (7) yields

$$\frac{\partial^2 f_A}{(\partial \theta_A)^2}(\theta_A, \theta_B) = 2l^2 \frac{T(\theta_A)}{(l^2 + \theta_B l^2 - 4 + \theta_A l^2 + \theta_A l^2 \theta_B)^4},$$

where

$$\begin{aligned}
T(\theta_A) = & -16 + \left(3\theta_B^2 + \theta_A + \theta_B^3 + 3\theta_B^2\theta_A + 3\theta_A\theta_B + 1 + \theta_A\theta_B^3 + 3\theta_B\right)l^5 \\
& + \left(16\theta_B + 14\theta_B^2 - 4\theta_B^2\theta_A - 2\theta_A\theta_B^3 + 6 - 2\theta_A\theta_B + 4\theta_B^3\right)l^4 \\
& + \left(-8\theta_B^2\theta_A - 8\theta_A - 16\theta_A\theta_B + 12\theta_B^2 + 24\theta_B + 12\right)l^3 \\
& + \left(-4\theta_B^2 - 8\theta_A\theta_B - 8\theta_A + 4\right)l^2 + (-16 - 16\theta_B)l
\end{aligned}$$

Somewhat tedious calculations reveal that $T(\theta_A) < 0$ if $l \in [-1/4, 3/5]$ and $\theta_A, \theta_B \in \Theta$, and thus $\frac{\partial^2 f_A}{(\partial \theta_A)^2}(\theta_A, \theta_B) < 0$, if $l \in [-1/4, 0) \cup (0, 3/5]$ and $\theta_A, \theta_B \in \Theta$. Therefore, the players' best replies concerning their overall biases are given by Eqs. (10). However, whether the players' best replies concerning their reciprocity biases are implicitly given by Eqs. (10) is still an open question. Differentiating Eqs. (7) with respect to α (respective β) by applying the chain rule leads to

$$\frac{\partial^2 f_1}{(\partial \alpha)^2}(\theta_1, \theta_2) = \frac{\partial^2 f_1}{(\partial \theta_1)^2}(\theta_1, \theta_2)(\gamma_1 - \gamma_2)^2,$$

and

$$\frac{\partial^2 f_2}{(\partial \beta)^2}(\theta_1, \theta_2) = \frac{\partial^2 f_2}{(\partial \theta_2)^2}(\theta_1, \theta_2)(\gamma_2 - \gamma_1)^2,$$

where $(\gamma_1 - \gamma_2)^2 > 0$ and $(\gamma_2 - \gamma_1)^2 > 0$ in any case, and $\frac{\partial^2 f_1}{(\partial \theta_1)^2}(\theta_1, \theta_2) < 0$, $\frac{\partial^2 f_2}{(\partial \theta_2)^2}(\theta_1, \theta_2) < 0$ if $l \in [-1/4, 0) \cup (0, 3/5]$. Then $\frac{\partial^2 f_1}{(\partial \alpha)^2}(\theta_1, \theta_2) < 0$ and $\frac{\partial^2 f_2}{(\partial \beta)^2}(\theta_1, \theta_2) < 0$ under the same constraint concerning l .

The next step is to show that the slope of the best reply functions of the asymmetric bias game are less than 1 in absolute value (cf. Lemma 1(ii)). We derive the two players' best reply functions concerning the reciprocity motive by calculating and solving the first order conditions of Eqs. (7) with respect to α (respective β), or equivalently, by plugging Eqs. (11) in Eqs. (10) and solving for the reciprocity norms. Accordingly, we reach

$$BR_1 = \alpha \left(\beta = \frac{\theta_2 - \gamma_1}{\gamma_2 - \gamma_1} \right) = \frac{(-l^2 + 2\gamma_2 l - 2l - \gamma_2 l^2)\theta_2 + 4\gamma_2 - \gamma_2 l^2 - l^2 + 2\gamma_2 l - 2l}{(l^2 - 2l)(\gamma_1 - \gamma_2)\theta_2 + (l^2 - 2l - 4)(\gamma_1 - \gamma_2)} \quad (13a)$$

$$BR_2 = \beta \left(\alpha = \frac{\theta_1 - \gamma_2}{\gamma_1 - \gamma_2} \right) = \frac{(\gamma_1 l^2 - 2\gamma_1 l + l^2 + 2l)\theta_1 + \gamma_1 l^2 - 2\gamma_1 l - 4\gamma_1 + l^2 + 2l}{(l^2 - 2l)(\gamma_1 - \gamma_2)\theta_1 + (l^2 - 2l - 4)(\gamma_1 - \gamma_2)}. \quad (13b)$$

The slopes of the best reply functions are given by

$$BR_1^s = \frac{d\alpha}{d\beta}(\beta) = -4 \frac{l(2+l)}{(4 + (\gamma_1\beta - \gamma_1 - \beta\gamma_2 - 1)l^2 + (-2\gamma_1\beta + 2 + 2\gamma_1 + 2\beta\gamma_2)l)^2} \quad (14a)$$

$$BR_2^s = \frac{d\beta}{d\alpha}(\alpha) = -4 \frac{l(2+l)}{(-4 + (\alpha\gamma_1 + 1 + \gamma_2 - \gamma_2\alpha)l^2 + (-2\alpha\gamma_1 - 2\gamma_2 - 2 + 2\gamma_2\alpha)l)^2}, \quad (14b)$$

where

$$\sup_{l \in [-1/4, 0) \cup (0, 3/5]} \left| -\left(4l^2 + 8l\right) \right| \approx 6, 24$$

and

$$\inf_{\substack{l \in [-1/4, 0) \cup (0, 3/5] \\ \gamma_1, \gamma_2 \in A \\ \alpha, \beta \in B}} |\text{den}(BR_1^s)| = \inf_{\substack{l \in [-1/4, 0) \cup (0, 3/5] \\ \gamma_1, \gamma_2 \in A \\ \alpha, \beta \in B}} |\text{den}(BR_2^s)| \approx 8, 265,$$

with $\text{den}(\circ)$ symbolizing the denominator of $\circ$. Since $6, 24 < 8, 265$ the slopes of the best reply functions of the two players are less than 1 in absolute value, so condition Lemma 1(ii) holds.

In conclusion, by Lemmata (1-3), the ex ante bias game $\Gamma_{\Gamma_f}^\gamma$ is dominance solvable with (α^*, β^*) as in Eq. (12) as the unique Nash equilibrium profile, and the unique profile that survives any dynamic regular and payoff-monotonic process with full support on the one-dimensional reciprocity space. $\square$

In order to analyze only situations where the evolutionary outcome of reciprocity lies between 0 and 1, we need to compute the following requirement as announced in Proposition 1.

Requirement. *The following relation is necessary to guarantee that $\alpha^*, \beta^* \in [0, 1]$.*

$$|\Phi(\gamma_1, \gamma_2, l) - (4l + 4)| \leq |l(l - 2)(\gamma_1 - \gamma_2)|. \quad (15)$$

Proof. Since $\alpha^* + \beta^* = 1$, it is sufficient to show that $\alpha^* \in [0, 1]$. Consider $0 \leq \alpha^* \leq 1$ with α^* as in Eq. (12), then we need to examine 2 cases:

Case 1: $l(l - 2)(\gamma_1 - \gamma_2) < 0$, i.e. $\text{sign}(l) = \text{sign}(\gamma_1 - \gamma_2)$, and

Case 2: $l(l - 2)(\gamma_1 - \gamma_2) > 0$, i.e. either $l < 0$ or $\gamma_1 < \gamma_2$.

Rearranging $0 \leq \alpha^* \leq 1$ given the first case leads to

$$4l + 4 + \underbrace{l(l - 2)(\gamma_1 - \gamma_2)}_{<0} \leq \Phi(\gamma_1, \gamma_2, l) \leq 4l + 4 - \underbrace{l(l - 2)(\gamma_1 - \gamma_2)}_{<0},$$

and the second case leads to

$$4l + 4 - \underbrace{l(l - 2)(\gamma_1 - \gamma_2)}_{>0} \leq \Phi(\gamma_1, \gamma_2, l) \leq 4l + 4 + \underbrace{l(l - 2)(\gamma_1 - \gamma_2)}_{>0}.$$

Subsuming both cases gives

$$4l + 4 - |l(l - 2)(\gamma_1 - \gamma_2)| \leq \Phi(\gamma_1, \gamma_2, l) \leq 4l + 4 + |l(l - 2)(\gamma_1 - \gamma_2)|,$$

and thus

$$|\Phi(\gamma_1, \gamma_2, l) - (4l + 4)| \leq |l(l - 2)(\gamma_1 - \gamma_2)|.$$

□

Solving for a parameter of this expression does not give much additional insight, instead, in the Appendix B we show some representative situations that conform to this requirement (see Figure 3 and Table 1 there).

Having proved our basic result, it is relatively simple to derive a benchmark finding where the evolutionary players are not able to feel reciprocity.²⁰ Let us first define a one-dimensional population as follows.

Definition 2. *A one-dimensional population here is a population as in the foregoing environment, but without reciprocity, and where the evolving trait is simply the intrinsic preference (the weight which is put on the opponents' material profit). That is, $\alpha = \beta = 1$ (cf. Definition 1), such that $\theta_A = 1 \cdot \gamma'_A + (1 - 1) \cdot \gamma'_B = \gamma'_A$ and $\theta_B = 1 \cdot \gamma'_B + (1 - 1) \cdot \gamma'_A = \gamma'_B$, where the “'” symbolizes “evolving” or “endogenized”. Also, in this setting, we allow for $\gamma'_A = \gamma'_B$.*

Then, we reach the following result.

Proposition 2. Consider the one-dimensional population described above with a strategic interdependence according to $l \in [-1/4, 0) \cup (0, 3/5]$ and a mutation space given by $\theta_A = \gamma'_A, \theta_B = \gamma'_B \in \Theta = [-1 + \epsilon, 1 - \epsilon]$. Then, any initial full-support distribution of biases converges in distribution towards the unit mass on $l/(2 - l)$, under any regular and payoff-monotonic dynamics.

Proof. As we have computed that $\frac{\partial^2 f_A}{(\partial \theta_A)^2}(\theta_A, \theta_B) < 0$ if $l \in [-1/4, 0) \cup (0, 3/5]$, it suffices to show that the slope of the best reply function is less than 1 in absolute value in this variant setting, since twice-differentiability and continuity, as also requested by Lemma 1, is obviously here. The best reply function of player A is

$$BR_A = \theta_A(\theta_B) = \operatorname{argmax}_{\theta_A} f_A(\theta_A, \theta_B) = -\frac{(\theta_B l + l + 2\theta_B + 2)l}{\theta_B l^2 + l^2 - 2\theta_B l - 2l - 4}. \quad (16)$$

²⁰Qualitatively, the same result appears in an example of Heifetz et al. (2007).

The slope of the best reply function is

$$BR_A^s = \frac{d\theta_A}{d\theta_B}(\theta_B) = \frac{4(l+2)l}{(\theta_B l^2 + l^2 - 2\theta_B l - 2l - 4)^2}. \quad (17)$$

Under assumptions $l \in [-1/4, 0) \cup (0, 3/5]$ and $\theta_A, \theta_B \in [-1 + \epsilon, 1 - \epsilon]$, we see that

$$\frac{dBR_A^s}{d\theta_B} = -8 \frac{(l+2)l(l^2 - 2l)}{(\theta_B l^2 + l^2 - 2\theta_B l - 2l - 4)^3} < 0.$$

Hence, Eq. (17) is decreasing in θ_B , and thus maximized at $\theta_B = -1 + \epsilon$ and minimized at $\theta_B = 1 - \epsilon$. With $l \in [-1/4, 0)$, Eq. (17) is negative and the maximum absolute value occurs at $\theta_B = 1 - \epsilon$, where $|BR_A(\theta_B = 1 - \epsilon)| = \left| 4 \frac{(l+2)l}{(-2l^2 + l^2\epsilon + 4l - 2l\epsilon + 4)^2} \right| < 1$. With $l \in (0, 3/5]$, Eq. (17) is positive and the maximum absolute value occurs at $\theta_B = -1 + \epsilon$, where $|BR_A(\theta_B = -1 + \epsilon)| = \left| 4 \frac{(l+2)l}{(l^2\epsilon - 2l\epsilon - 4)^2} \right| < 1$. Since the variant game is dominance solvable, we find the outcome $l/(2-l)$ by equalizing and solving for the biases with respect to Eq. (16), i.e. solving for θ_A in $\theta_A(\theta_B = \theta_A)$. $\square$

So, what is the significance of these results? By fielding this question, one should bear in mind that although the constraints of the equilibrium, which are determined by the game dynamical aspects and the model parameters, appear somewhat restricted, all assumption are intuitively plausible and of sufficient general character. The game dynamics allow for different initial populations to develop in the wide field of regularity and payoff monotonicity only provided that the population describes a compact interval in the line of reciprocity. Likewise, the payoff function which defines reproductive success in the society stands for a wide variety of different strategic games.

There are several observations which we can make easily. To start with, note that both players' equilibrium reciprocity values sum up to 1, i.e. $\alpha^* + \beta^* = 1$. This fact guarantees that both players' regard for the opponents' payoff is identical.

Corollary 1. $\theta_1^* = \theta_2^*$.

Proof. Since $\alpha^* + \beta^* = 1$, we have $\theta_1^* = \underbrace{\alpha^*}_{=1-\beta^*} \cdot \gamma_1 + \underbrace{\beta^*}_{=1-\alpha^*} \cdot \gamma_2 = \theta_2^*$. $\square$

This result is not surprising because asymmetry of the equilibrium payoffs emerges not on the level of the players' overall concern but only with respect to the intrinsic preferences of the players. The next observation regards a comparing of our basic result with the outcome in the one-dimensional population model.

Corollary 2. *The two-dimensional population model may develop a different overall concern than the one-dimensional population model does.*

Proof. Let us write down only one simple example. Consider $l = -0,25$, $\gamma_1 = -0,8$, and $\gamma_2 = 0,4$. Then, the two-dimensional population develops an outcome that is approximately $\alpha^* \cdot \gamma_1 + \beta^* \cdot \gamma_2 \approx -0,08$ and the outcome of the one-dimensional population model corresponds to $\frac{l}{2-l} = -0, \bar{1}$. $\square$

However, the fact that the strategic environment determines the sign of the overall concern is an observation which is of general character.

Corollary 3. *The players show a negative overall concern if strategic substitutes are present, i.e. $l < 0 \Rightarrow \theta_1^* = \theta_2^* < 0$. If the underlying game exhibits strategic complements, then, the players' value their opponents' payoff positively, i.e. $l > 0 \Rightarrow \theta_1^* = \theta_2^* > 0$.*

This result is reminiscent of the pioneering work of [Bester and Güth \(1998\)](#) where strategic complements leads to altruism and strategic substitutes to selfishness.²¹ The following observation regards the reciprocity motive and its conclusion holds universally for the case of strategic substitutes, i.e. $l \in [-1/4, 0)$, and partly for the case of strategic complements ($l \in (0, 3/5]$).

Corollary 4. *There exists a threshold $\zeta(l) \in A$ such that*

$$\frac{\partial \alpha^*(\gamma_1, \gamma_2, l)}{\partial \gamma_2} < 0, \text{ for } \gamma_1 < \zeta(l), \quad (18)$$

and

$$\frac{\partial \alpha^*(\gamma_1, \gamma_2, l)}{\partial \gamma_2} > 0, \text{ for } \gamma_1 > \zeta(l), \quad (19)$$

and for the second player likewise. A rough conclusion of this observation works as follows. If player A's intrinsic lies below the threshold, then the will to reciprocate to player B increases as player B gets more nasty. Likewise, if player A's intrinsic lies above the threshold, then the will to reciprocate to player B increases as player B gets more nice. Furthermore, the threshold is increasing in the strategic setting l .

As the fractions of Ineq. (18) and Ineq. (19) are no continuous functions for the parameter constraints when strategies are complements, the conclusion derived from these inequations is limited for $l \in (0, 3/5]$. In particular, for strategic complements and some values of $\gamma_1 \in A$, denote $\hat{\gamma}_1$, there is a critical point for γ_2 such that $\alpha^*(\hat{\gamma}_1, \gamma_2 - \delta, l \in (0, 3/5]) < \alpha^*(\hat{\gamma}_1, \gamma_2, l \in (0, 3/5]) > \alpha^*(\hat{\gamma}_1, \gamma_2 + \delta, l \in (0, 3/5])$, for positive δ , and for the second player likewise.

²¹However, the finding under strategic substitutes is restricted there, which is due to Bester and Güth's presumption of a non-negative preference space, and extends to spitefulness if the zero-barrier is abrogated (cf. [Bolte, 2000](#); [Possajennikov, 2000](#)).

4 Conclusion

Building upon the assumption that individuals adjust their actions to achieve higher subjective utility, while dynamical pressures change the composition of reciprocal preferences in the population according to the players' objective gains, this study provides a prognosis concerning the emergence of strong reciprocity in a wide class of strategic interaction. The specific conception of reciprocity is defined by the players tendency to response to the perceived intrinsic attitudes of others. The basic finding is that a high degree of flexibility (in the reciprocal sense) pays off. In our setting, it is the strategic environment and the specific other players' type which determine the players' behavior and only marginally their usual, exogenous given, intrinsic attitudes. The unique dominance solvable profile of reciprocity in our asymmetric setting, and hence the only survivor under any regular and payoff monotonic selection process, motivates an altruistically inclined player to behave spitefully if strategies are substitutes and the opponent player is intrinsically spiteful. Conversely, the reciprocity norm of a spiteful player motivates him to show altruistic behavior if strategies are complements and the opponent player is intrinsically altruistic. In this regard, our study provides an economic and cultural explanation to the question why people often show a different behavior than they usually would. The study further substantiates related work which shows that the strategic environment determines the players equilibrium behavior in one-dimensional preference models. In particular, the fact that strategic substitutes lead to a negative emotion and strategic complements to a positive emotion is once more confirmed, even in the new setting of two preference dimensions.

A further conclusion of our work is that if player A is relatively spiteful, then player A's tendency to reciprocity is higher the nastier player B—if player A is relatively nice, then player A's tendency to reciprocity is higher the nicer player B. However, this conclusion is restricted in the sense that it holds universally only for the case of strategic substitutes.

One can think of several extensions of our basic model. Since preference biases act like commitment devices which form the other players' equilibrium strategies to a certain extent, it would be interesting to explore whether the qualitative results maintain in cases where the players do not recognize the other players' types perfectly. For instance, one can think of situations where the players' types are observed with some noise because they do not learn the types of each other, or where some players intentionally signal a wrong disposition in order to benefit from the resulting strategic effect. It would further be interesting whether our results can be identified in experimental studies—an admittedly subtle task since the individual usual types have to be ascertained before a norm of reciprocity can be examined.

Appendix A. A generic class of evolutionary dynamics

To analyze the evolutionary viability of reciprocity, we build upon the generic class of selection dynamics as proposed by [Heifetz et al. \(2007\)](#). This section is intended to sketch the distinctive attributes and advantages of this class. To this end, imagine the following.

At each instant in time $t \geq 0$, two players are randomly drawn from a continuum population to play a certain game. In particular, the population is characterized by the distribution $G_t \in \Delta(\Theta)$, where $\Delta(\Theta)$ is the set of Borel probability distributions over the compact space $\Theta = [\underline{\theta}, \bar{\theta}]$. The population evolves over time in the space of $\Delta(\Theta)$ according to the following differential equation.

$$\dot{G}_t(S) = \int_S g(\theta, G_t) dG_t(\theta), \quad S \subseteq \Theta \text{ Borel measurable}, \quad (\text{A.1})$$

where $g : \Theta \times \Delta(\Theta) \rightarrow \mathbb{R}$ is a continuous growth-rate function. The following definition further specifies the dynamics.

Definition 3. *The continuous growth-rate function $g : \Theta \times \Delta(\Theta) \rightarrow \mathbb{R}$ is payoff-monotonic and regular if for every $G \in \Delta(\Theta)$, the following conditions hold:*

- *A higher average fitness corresponds to a higher growth-rate, or formally,*

$$\int f(\theta, \hat{\theta}) dG_t(\hat{\theta}) > \int f(\tilde{\theta}, \hat{\theta}) dG_t(\hat{\theta}) \iff g(\theta, G_t) > g(\tilde{\theta}, G_t). \quad (\text{A.2})$$

- *G_t is a probability distribution for every t ,*

$$\int_{\Theta} g(\theta, G) dG(\theta) = 0. \quad (\text{A.3})$$

- *g can be extended to the domain $\Theta \times X$, where X is the set of signed Borel measures with variational norm smaller than 2, such that g is uniformly bounded and Lipschitz continuous on $\Theta \times X$. Formally,*

$$\begin{aligned} \sup_{\theta \in \Theta} |g(\theta, G_t)| &< \infty \\ \sup_{\theta \in \Theta} |g(\theta, G_t) - g(\theta, \tilde{G}_t)| &< K \|G_t - \tilde{G}_t\|, \quad G_t, \tilde{G}_t \in X, \end{aligned} \quad (\text{A.4})$$

for some constant K , where $\|G\| = \sup_{|h| \leq 1} |\int_{\Theta} h dG|$ is the variational norm on signed measures.

[Oechssler and Riedel \(2001, Lemma 3\)](#) show that regularity of g guarantees that the mapping $G \rightarrow \int_{\Theta} g(\cdot, G) dG$ is bounded and Lipschitz continuous in the variational norm,

which implies that the differential Eq. (A.1) has a unique solution for any initial distribution G_0 .²²

The dynamics defined here formalize the simple idea that only individuals who play well in the population increase while individuals who play badly decrease. As mentioned in the introduction, the underlying evolving process may rely on a biological level or on a cultural level. Accordingly, more successful types have more descendantes who carry the genes of their parents, or more successful types are more likely to be adopted under a cultural process of education or imitation from role-models. Alternatively, [Heifetz et al. \(2007\)](#) mention that the same mathematical structure is compatible with the idea that successful individuals have more influence on the dynamic process since they appear more often in economic interactions, and so, are more likely to be reproduced.

Dealing with the concept of dominance solvability requires a deeper understanding of the concept of domination and some additional notations. To begin with, we say that θ' is dominated by θ whenever $f(\theta, \theta'') > f(\theta', \theta'')$ for every $\theta'' \in \Theta$. Then, let D_1 denote the set of types θ' which are dominated by some $\theta \in \Theta$, and U_1 is the set of undominated types, i.e. $U_1 = \Theta \setminus D_1$. Further, D_n is the set of dominated types after at most n iterations and the set of undominated types is accordingly $U_n = \Theta \setminus D_n$. Then, $\theta' \in U_n$ is dominated in iteration $n + 1$ by $\theta \in U_n$ if $f(\theta, \theta'') > f(\theta', \theta'')$ for every $\theta'' \in U_n$. We say that θ' is *serially dominated* if it is dominated after some number of iterations. Under regular and payoff monotonic dynamics any serially dominated types are extinct in the limiting population. A game is dominance solvable if there is a unique type that is not serially dominated. The analyzes of reciprocity in this paper is based on this class of selection dynamics and on the following theorem of [Heifetz et al. \(2007\)](#) which extends results from [Samuelson and Zhang \(1992\)](#), where the population evolves according to matrix games, to continuous strategy spaces.

Theorem. ([Heifetz et al. \(2007\)](#)) *Consider a symmetric two-player game with strategy space $\Theta = [\underline{\theta}, \bar{\theta}] \subset \mathbb{R}$, a continuous payoff function $f : \Theta \times \Theta \rightarrow \mathbb{R}$, and a regular, payoff-monotonic growth-rate function $g : \Theta \times \Delta(\Theta) \rightarrow \mathbb{R}$. Moreover, assume that the*

²²To see that the replicator dynamics forms a special case of the generic dynamics determined by Eq. (A.1), consider a growth-rate function which evolves according to the subtraction of the population average success from the success of a single type (the difference is sometimes called the *excess payoff*). Formally, G_t evolves according to

$$\dot{G}_t(S) = \int_S [f(\theta_i, G_t) - f(G_t, G_t)] dG_t(\theta_i), \quad S \subseteq \Theta,$$

where

$$f(\theta_i, G_t) = \int_{\Theta} f_i(\theta_i, \theta_j) dG_t(\theta_j)$$

is the expected success of individual i played against a randomly chosen individual j . And, the population average success is

$$f(G_t, G_t) = \int_{\Theta} \int_{\Theta} f_i(\theta_i, \theta_j) dG_t(\theta_j) dG_t(\theta_i).$$

population G has initially full support over the compact space Θ and evolves according to the differential equation as defined by Eq. (A.1). Then, types θ which are serially dominated are asymptotically weeded out, i.e. they have a neighborhood $V \ni \theta$ for which $\lim_{t \rightarrow \infty} G_t(V) = 0$. In particular, when the game is dominance solvable to equilibrium θ^* , then G_t converges in distribution to a unit mass at θ^* .

Appendix B. Figure and Table

Figure 3 presents the possible sets of γ_1 and γ_2 in 5 different strategic situations l : $l = -0,25$, $l = -0,05$, $l = 0,05$, $l = 0,25$, and $l = 0,5$. The graphics show that we analyze situations where the intrinsic preferences of the two populations are predominantly different, i.e. $\text{sign}(\gamma_1)$ is predominantly different from $\text{sign}(\gamma_2)$. The graphics also show that in the case of strategic substitutes $l < 0$, we do not consider populations with positive γ_1 and positive γ_2 ; likewise, in the case of strategic complements, we do not consider cases where γ_1 and γ_2 are both negative.

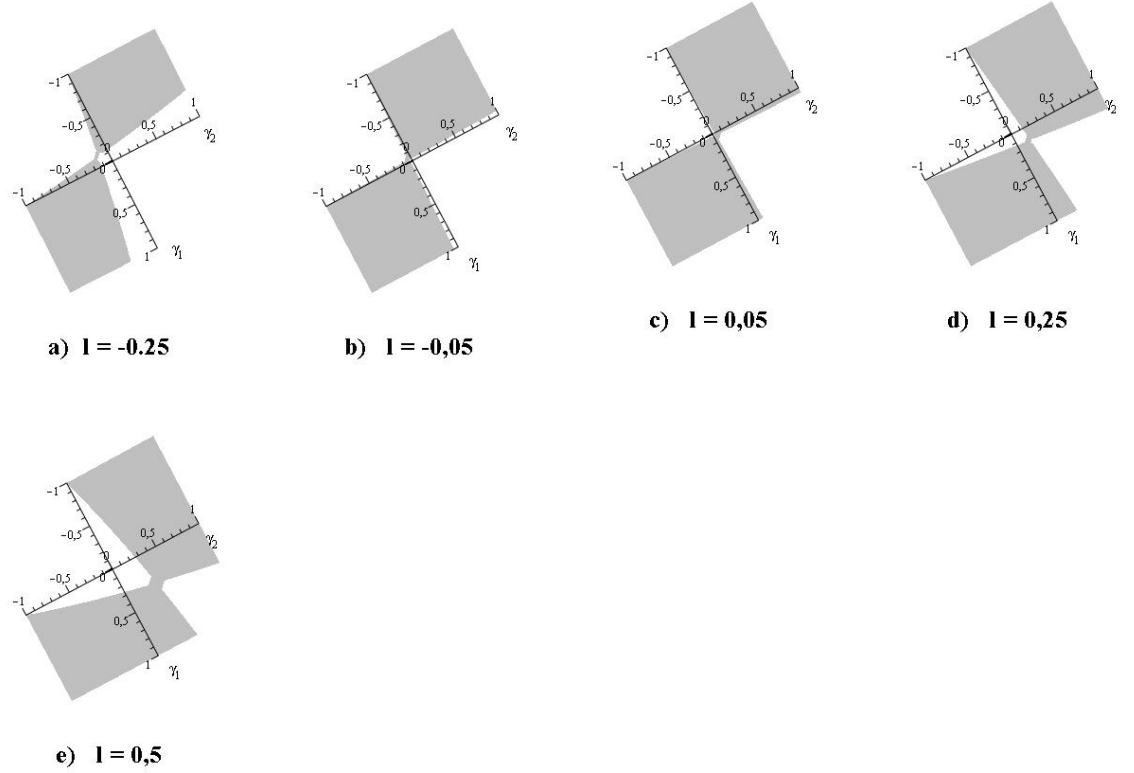


Figure 3: γ_1, γ_2 -sets.

Table 1 shows some dominance solvable results with respect to different game parameters. For expositional clarity we focus on the same “representative” environments as in Figure 3: we assume 5 different cases which examine the strategic environment, i.e. $l = -0,25$, $l = -0,05$, $l = 0,05$, $l = 0,25$, and $l = 0,5$, each with a two population setting of different intrinsic preferences. The table basically shows that a dominance solvable reciprocity trait is apt to reverse the players intrinsic attitude. This finding holds in settings with strategic substitutes ($l < 0$) as well as in those with strategic complements ($l > 0$). For example take $l = -0,25$, i.e. strategic substitutes, then if the γ_1 -player is moderate altruistic ($\gamma_1 = 0,3$) and the opponent player is of type $\gamma_2 = -0,3$ then the dynamics drive the γ_1 -player to a relatively reciprocal norm $(0,25)$ such that the overall concern becomes negative. The fact that the sign of the strategic environment determines the sign of the overall concern can also be observed in the table.

Exogeneous game parameters			Dominance solvable traits (approximate values)		
Strategic setting	Intrinsic preferences		Reciprocity		Overall concern
l	γ_1	γ_2	α^*	β^*	$\theta_1^* = \theta_2^*$
$-0,25$	$-0,9$	$0,1$	$0,125$	$0,875$	$-0,025$
$-0,25$	$0,3$	$-0,3$	$0,25$	$0,75$	$-0,15$
$-0,05$	$-0,6$	$0,3$	$0,35$	$0,65$	$-0,015$
$-0,05$	$0,6$	$-0,1$	$0,087$	$0,913$	$-0,039$
$0,05$	$0,7$	$-0,2$	$0,262$	$0,738$	$0,036$
$0,05$	$-0,7$	$0,5$	$0,4$	$0,6$	$0,02$
$0,25$	$0,4$	$-0,5$	$0,67$	$0,33$	$0,103$
$0,25$	$-0,3$	$0,2$	$0,19$	$0,81$	$0,105$
$0,5$	$0,6$	$0,1$	$0,478$	$0,522$	$0,339$
$0,5$	$-0,7$	$0,3$	$0,166$	$0,834$	$0,134$

Table 1: Results in different situations.

References

- ANDREONI, J., P. BROWN, AND L. VESTERLUND (2002): “What Produces Fairness? Some Experimental Evidence,” *Games and Economic Behavior*, 40, 1–24.
- ANDREONI, J. AND J. MILLER (2002): “Giving According to GARP: An Experimental Test of the Consistency of Preferences for Altruism,” *Econometrica*, 70, 737–753.
- BESTER, H. AND W. GÜTH (1998): “Is Altruism Evolutionary Stable?” *Journal of Economic Behavior and Organization*, 34, 193–209.
- BOLLE, F. (2000): “Is Altruism Evolutionarily Stable? And Envy and Malevolence? - Remarks on Bester and Güth,” *Journal of Economic Behavior and Organization*, 42, 131–133.
- BOLTON, G. AND A. OCKENFELS (2000): “ERC: A Theory of Equity, Reciprocity and Competition,” *American Economic Review*, 90, 166–193.
- BULOW, J., J. GENEAKOPOLOS, AND P. KLEMPERER (1985): “Multimarket Oligopoly: Strategic Substitutes and Complements,” *Journal of Political Economy*, 93, 488–511.
- CAMERER, C. (2003): *Behavioral Game Theory: Experiments on Strategic Interaction*, Princeton University Press.
- CHARNESS, G. AND M. RABIN (2002): “Understanding Social Preferences with Simple Tests,” *Quarterly Journal of Economics*, 117, 817–869.
- DUFWENBERG, M. AND G. KIRCHSTEIGER (2004): “A Theory of Sequential Reciprocity,” *Games and Economic Behavior*, 47, 268–298.
- FALK, A. AND U. FISCHBACHER (2006): “A Theory of Reciprocity,” *Games and Economic Behavior*, 54, 293–315.
- FEHR, E. AND U. FISCHBACHER (2002): “Why Social Preferences Matter—the Impact of Non-selfish Motives on Competition, Cooperation, and Incentives,” *The Economic Journal*, 112, C1–C33.
- FEHR, E. AND K. SCHMIDT (1999): “A Theory of Fairness, Competition, and Cooperation,” *Quarterly Journal of Economics*, 114, 817–868.
- FRANK, R. (1987): “If Homo Economicus Could Choose His Own Utility Function, Would He Want One with a Conscience?” *American Economic Review*, 77, 593–604.
- (1988): *Passions within Reason: The strategic Role of the Emotions*, W.W. Norton, New York.
- GEANAKOPOLOS, J., D. PEARCE, AND E. STACCHETTI (1989): “Psychological Games and Sequential Rationality,” *Games and Economic Behavior*, 1, 60–79.
- GINTIS, H. (2000): “Strong Reciprocity and Human Sociality,” *Journal of Theoretical Biology*, 206, 169–179.
- GÜTH, W., R. SCHMITTBERGER, AND B. SCHWARZE (1982): “An Experimental Analysis of Ultimatum Bargaining,” *Journal of Economic Behavior and Organisation*, 3, 367–388.
- GÜTH, W. AND M. YAARI (1992): “Explaining Reciprocal Behavior in Simple Strategic Games: An Evolutionary Approach,” in *Explaining Forces and Changes: Approaches to Evolutionary Economics*, ed. by U. Witt., University of Michigan Press.
- HARRISON, R. AND M. VILLENA (2008): “On the Evolution of Reciprocal Behavior: A Game Theoretic Approach,” Working Paper.
- HEIFETZ, A. AND E. SEGEV (2004): “The Evolutionary Role of Toughness in Bargaining,” *Games and Economic Behaviour*, 49, 117–134.
- HEIFETZ, A., C. SHANNON, AND Y. SPIEGEL (2007): “The Dynamic Evolution of Preferences,” *Economic Theory*, 32, 251–286.

- HOFBAUER, J., J. OECHSSLER, AND F. RIEDEL (2009): “Brown-von Neumann-Nash Dynamics: The Continuous Strategy Case,” *Games and Economic Behavior*, 65, 406–429.
- HOFFMANN, E., K. MCCABE, AND V. SMITH (1998): “Behavioral Foundations of Reciprocity: Experimental Economics and Evolutionary Psychology,” *Economic Inquiry*, 36, 335–352.
- HUCK, S. AND J. OECHSSLER (1999): “The Indirect Evolutionary Approach to Explaining Fair Allocations,” *Games and Economic Behavior*, 28, 13–24.
- LEVINE, J. (1998): “Modeling Altruism and Spitefulness in Experiments,” *Review of Economic Dynamics*, 1, 593–622.
- MAYNARD-SMITH, J. AND G. PRICE (1973): “The Logic of Animal Conflicts,” *Nature*, 246, 15–18.
- MOULIN, H. (1984): “Dominance Solvability and Cournot Stability,” *Mathematical Social Sciences*, 7, 83–102.
- NOWAK, M. AND K. SIGMUND (2005): “Evolution of Indirect Reciprocity,” *Nature*, 437, 1291–1298.
- OECHSSLER, J. AND F. RIEDEL (2001): “Evolutionary Dynamics on Infinite Strategy Spaces,” *Economic Theory*, 17, 141–162.
- POSSAJENNIKOV, A. (2000): “On the Evolutionary Stability of Altruistic and Spiteful Preferences,” *Journal of Economic Behavior and Organisation*, 42, 125–129.
- RABIN, M. (1993): “Incorporating Fairness into Game Theory and Economics,” *American Economic Review*, 83, 1281–1302.
- SAMUELSON, L. AND J. ZHANG (1992): “Evolutionary Stability in Asymmetric Games,” *Journal of Economic Theory*, 57, 363–391.
- SCHELLING, T. (1960): *The Strategy of Conflict*, Harvard University Press, Cambridge.
- SELTEN, R. (1980): “A Note on Evolutionary Stable Strategies in Asymmetric Animal Conflicts,” *Journal of Theoretical Biology*, 84, 93–101.
- (1991): “Evolution, Learning, and Economic Behavior,” *Games and Economic Behavior*, 3, 3–24.
- SETHI, R. AND E. SOMANTHAN (2001): “Preference Evolution and Reciprocity,” *Journal of Economic Theory*, 97, 273–297.
- (2003): “Understanding Reciprocity,” *Journal of Economic Behavior and Organization*, 50, 1–27.
- SOBEL, J. (2005): “Interdependent Preferences and Reciprocity,” *Journal of Economic Literature*, 43, 396–440.
- TRIVERS, R. (1971): “The Evolution of Reciprocal Altruism,” *Quarterly Review of Biology*, 46, 35–58.
- WEIBULL, J. (1995): *Evolutionary Game Theory*, The MIT Press, Cambridge, MA.