

Küpper, Hans-Ulrich; Sandner, Kai

Working Paper

Differences in Social Preferences - Are They Profitable for the Firm?

Discussion Paper, No. 2008-03

Provided in Cooperation with:

University of Munich, Munich School of Management

Suggested Citation: Küpper, Hans-Ulrich; Sandner, Kai (2008) : Differences in Social Preferences - Are They Profitable for the Firm?, Discussion Paper, No. 2008-03, Ludwig-Maximilians-Universität München, Fakultät für Betriebswirtschaft, München, <https://doi.org/10.5282/ubm/epub.2122>

This Version is available at:

<https://hdl.handle.net/10419/104493>

Standard-Nutzungsbedingungen:

Die Dokumente auf EconStor dürfen zu eigenen wissenschaftlichen Zwecken und zum Privatgebrauch gespeichert und kopiert werden.

Sie dürfen die Dokumente nicht für öffentliche oder kommerzielle Zwecke vervielfältigen, öffentlich ausstellen, öffentlich zugänglich machen, vertreiben oder anderweitig nutzen.

Sofern die Verfasser die Dokumente unter Open-Content-Lizenzen (insbesondere CC-Lizenzen) zur Verfügung gestellt haben sollten, gelten abweichend von diesen Nutzungsbedingungen die in der dort genannten Lizenz gewährten Nutzungsrechte.

Terms of use:

Documents in EconStor may be saved and copied for your personal and scholarly purposes.

You are not to copy documents for public or commercial purposes, to exhibit the documents publicly, to make them publicly available on the internet, or to distribute or otherwise use the documents in public.

If the documents have been made available under an Open Content Licence (especially Creative Commons Licences), you may exercise further usage rights as specified in the indicated licence.

Differences in Social Preferences: Are They Profitable for the Firm?

Hans-Ulrich Küpper und Kai Sandner

Discussion Paper 2008-3

November 2008



Munich School of Management

University of Munich

Fakultät für Betriebswirtschaft

Ludwig-Maximilians-Universität München

Online at <http://epub.ub.uni-muenchen.de/>

Differences in Social Preferences: Are They Profitable for the Firm?

Hans-Ulrich Küpper

*Institute of Production Management and Management Accounting
Munich School of Management
University of Munich
Ludwigstrasse 28, 80539 Munich, Germany
kuepper@bwl.lmu.de*

Kai Sandner

*Institute of Production Management and Management Accounting
Munich School of Management
University of Munich
Ludwigstrasse 28, 80539 Munich, Germany
sandner@bwl.lmu.de*

This paper analyzes the impact of heterogeneous (social) preferences on the weighting and combination of incentive performance measures as well as on a firm's profitability within a principal-agent framework. Previous literature had failed to recognize heterogeneity effects. We consider rivalry, pure self-interest and altruism as extreme forms of such preferences within the spectrum of possible alternatives, and show that firm profits are maximized when differentiation among agents is maximized with respect to individual (social) preferences. In order to realize these gains in profitability, it is necessary that a firm directs principals to reallocate participation in performance measures so that competitive agents are privileged over altruistic agents. By modeling the need to incorporate heterogeneity of agents we provide insights that differences in social preferences can be managed to improve wage compensation and other business administration deliberations within a decentralized firm.

1. Introduction

If we try to define the concept of 'firm' as simply as possible, we could say that firms are structures where several people work together in order to achieve common goals. Humans are social beings, and as such they may care, or not care, for one another and inevitably bring with them differences in views and tastes. With that in mind, it is clear that not all employees get along with each other equally well, especially in places where interaction is frequent and personal, such as firms. Some employees are more productive in certain environments than in others. This is partly a result of the motivational skills of direct leaders; however, differences in personal social preferences also seem to have a major influence on employee productivity.

Using a principal-agent framework within a decentralized organization two questions are examined: (1) what factors determine the impact that differences in social preferences have on profitability? (2) How can we compose a group of workers optimally in order to guarantee the best possible performance from a superordinate's point of view? To answer these questions, we will first show how a principal reacts to his or her agents' homogeneous and heterogeneous social preferences of varying strength when incentives are provided. Following that, we will present formally derived statements about the consequences of social preferences on firm profitability.

Previous works have failed to recognize the heterogeneity of agents in their analyses. When studying organizational behavior, economics experts have so far typically concentrated on analyzing the structure of the optimal reward system under the assumption that employees behave in a purely selfish manner.¹ Since those studies do not consider differences in preferences, it has been difficult to draw conclusions about the best way of composing a team. However, studies from different fields within the behavioral sciences like psychology², neuroscience³ and

experimental decision theory⁴ suggest that, in many cases, the economic decisions that individuals make are also determined by social preferences such as altruism, inequity aversion and reciprocity.⁵

Generally, preferences are designated as social “if a person not only cares about the material resources allocated to her but also cares about the material resources allocated to relevant reference agents.”⁶ On the basis of this definition, various types of motivational structures are distinguished in the economics literature,⁷ the main criterion being the effect that changes in another person’s payoff has on the individual’s own utility. In this paper, our theoretical framework does not permit us to analyze all relevant types of preferences. We have chosen to focus on pure self-interest, rivalry and altruism, although reciprocity and inequity aversion are also examined to some extent in the previous literature, especially in the context of experimental decision theory. An important reason for our choice is that we are interested in the possible effects of differences between the preferences of several agents on the principal’s provision of incentives and on firm profitability. Therefore, social preferences that aim at a minimum of differences e.g. in wages, like reciprocity and inequity aversion, cannot be a suitable starting point for such an investigation. By contrast, rivalry, pure self-interest and altruism represent three types of preferences within the spectrum of possible types, where at the one extreme, in the case of rivalry, the other agent’s payoff is always evaluated negatively, and at the other extreme, in the case of altruism, it is always evaluated positively, independently of the overall distribution of wages.

Focusing on those three types allows us to draw intuitive conclusions about the either complementary or contradictory relation between a principal’s counteractions (necessary for exploiting each of his or her single agents’ social preferences) and, as a consequence, about the impact that such measures have on

firm profitability. Furthermore, in reality, rivalry seems to play an important role in many business areas, as can be seen from the ongoing discussion on wage justice in Germany⁸ as well as in Northern America⁹, but also in other fields of interest such as sports. Altruism, however, has also been shown to have a strong effect on human behavior under various circumstances.¹⁰

As we will see below, rivalry and altruism can be modeled in comparable terms,¹¹ each time including pure self-interest as a limit in case we take the particular social preference parameters in our mathematical formulations to zero. This enables us to employ the same analytical framework for studying differences in intensity among homogeneous as well as among heterogeneous types of (social) preferences. Therefore, in this paper we analyze one formal class of comparable preferences, of which three extreme types have been included. By contrast, preferences like reciprocity and inequity aversion, which are based on fairness, would require a different modeling approach.

Our results highlight the importance of variety in the agents' motivational structures. In the following, we will show that a principal can use this variety to increase his or her own profits. As a prerequisite, he or she must shift the weightings of performance measures between agents in a way that the (more) competitive agent gains at the expense of the other agent. This scheme can be seen as an optimal way of reacting to the agents' social preferences. If the monetarily discriminated agent views such measures as positive, this translates into higher profits for the firm, i.e. the less "disadvantaged" agents see the consequences of such measures as negative, the higher the firm's profits, which implies that principals can often make use of differences in motivational structures, but cannot profit from equally strong homogeneous social preferences. This contradicts the common intuition that altruistic agents, who do not begrudge others their prosperity, make the largest contributions,

and are therefore the most important members of a team. Instead, as will become evident, competitive persons are equally significant as part of a mixture of different characters.

The remainder of the paper is organized as follows: Section 2 relates our work to previous research in the field. Section 3 outlines the framework of our analysis and describes the theoretical model we use here with its basic assumptions. Sections 4 and 5 contain our main results. First, the consequences of taking into account different social preferences when designing the structure of the optimal wage compensation system are examined in Section 4. Then, on the basis of our results, in Section 5 we consider the impact of heterogeneous (social) preferences on firm profitability. Section 6 explores the implications of our theoretical analysis and highlights opportunities for further research.

2. Relation to the Literature

Previous research on social preferences in principal-agent models distinguishes between models that concentrate on vertical comparisons of an agent with his or her principal¹² and models that concentrate on horizontal comparisons between agents on the same horizontal layer.¹³ In the latter case, authors usually focus on the optimal design of the wage compensation system.¹⁴ The question of whether the principal can make use of his or her agents' social preferences is explicitly addressed in only a few papers, which do not offer clear-cut results.

Rey Biel (2007) examines the impact of inequity aversion and shows that in a situation of team production without any uncertainty, welfare comparisons between agents have an intrinsic motivation effect on individual efforts. Itoh (2004) confirms these results and generalizes them to include competitive preferences in a model which incorporates uncertainty but neglects technological dependences. He draws

the conclusion that social preferences among agents on the same hierarchical layer are generally beneficial from a firm's point of view. Both papers impose limited liability constraints. In this spirit, Bartling and von Siemens (2006) claim that among employees, envy can only have a cost-reducing effect in cases of limited liability.

By contrast, Dierkes and Harreiter (2006) show that a principal can make use of his or her agents' rivalry also in cases where there is no limited liability; however, Bartling and von Siemens (2005, 2006) point out that in such cases, envy and a preference for inequity aversion may increase agency costs and thus harm a firm's profits. Neilson and Stowe (2008) suggest that, despite an intrinsic motivation effect, in cases where inequity-averse agents are concerned, the principal's inequity premium increases with the inequity aversion of his or her agents.¹⁵ Similarly, Demougin, Fluet and Helm (2006) deduce that in a two-task environment, where only one agent's efforts are observable, inequity aversion unambiguously decreases total output and hence labor productivity, although their model includes the limited liability constraints.

As this brief overview shows, the available literature offers no clear understanding of how different social preferences impact profitability. What's more, all cited papers take for granted that agents do not differ with respect to types of (social) preferences while the issue of altruism is not adequately discussed. In actual business life, however, it is more realistic to assume that the members of a team may have different behavioral traits. In our study, we shall try to provide a thorough understanding of the factors that determine the impact of diverse types of social preferences on profitability. To achieve this, the underlying theoretical framework must allow the explicit computation of utilities and, therefore, a precise solution. Previously developed analytical models are often formulated in a very general manner and hence cannot be solved for the principal's decision parameters that

specify his or her contract offer¹⁶ or only allow remuneration to be assessed in dependence of discrete outcomes that are affected by agents' discrete efforts as well as by which discrete and stochastic external conditions apply.¹⁷ In the latter case, the explicit calculation of utility values is in fact possible, but only under additional restrictive assumptions, such as that agents have identical (social) preferences.¹⁸ Therefore, we employ an exactly solvable principal-multiagent linear-exponential-normal (LEN) model¹⁹ and extend it by incorporating heterogeneous social preferences.

Concerning the optimal provision of incentives, earlier studies have shown that preferences for inequity aversion (Itoh 2004; Bartling and von Siemens 2005) as well as envy or rivalry (Bartling and von Siemens 2006; Dierkes and Harreiter 2006) lead principals to provide their agents with more equitable, so-called "flat wage" contracts. Itoh (2004) also compares the advantages of team-based compensation and relative performance evaluation under the assumption that agents have identical (social) preferences. Although this is not our primary objective, in our paper we explore this point further, by investigating how to weight and combine performance measures in the presence of various social preferences as well as stochastic dependences.²⁰ In contrast to most existing studies,²¹ we also allow for differences in the agents' other personality traits, apart from their motivational structures.²² By doing so, we are able to deliver new, in-depth insights into the overlapping functions of the principal's share rates, as well as into their motivational potential.

3. The Model

Conceptual Framework of Analysis

The impact of social preferences on the optimal structure of the principal's system of incentives and firm profits may depend on both external and internal determinants.

Organizational and environmental conditions, as well as conditions of production, are of particular importance. The relevance of organizational conditions derives from the assumption that agents belong to different decentralized divisions. Within a firm, arguably an employee's most important connections to his or her superior follow a vertical direction, while those to his or her colleagues on the same hierarchical level follow a horizontal direction (see Figure 1). Both types of relationships may impact an employee's performance. In this paper, we are mainly interested in how the heterogeneous preferences among agents affect firm performance, so we will only consider the psychological interdependences between pairs of horizontally aligned agents on the same hierarchical layer.

Insert Figure 1 here

Figure 1: Directions of psychological relationships as a result of social preferences

Technological dependences in the production process arise if an agent's performance measure is affected not only by his or her own activities, but also by those of agents in other decentralized divisions (and vice versa). Environmental stochastic dependences are present if the profits of two decentralized divisions depend on correlated error terms. For example, such correlations can result from external market conditions or the general business cycle. Therefore, the influences of different preferences in a decentralized organization can be analyzed in cases of internal technological and/or external stochastic dependence as well as independence (see Figure 2).

Insert Figure 2 here

Figure 2: Determinants of the research problem

Space does not allow us to explore fully the influence of technological production dependences, so in this paper we shall focus on the interrelation between stochastic

and psychological dependences.²³ However, elsewhere it has been shown that our main results about the impact of social preferences on firm profitability hold also in the case of technological dependences.²⁴

The structure and impact of preference interdependences are closely linked to an agent's personal traits. Social preferences typically concern one or several other persons. In our model we will consider pairs of agents on the same hierarchical level (Figure 1). The two agents may have competitive, selfish or altruistic preferences. Thus, three homogeneous and six heterogeneous combinations can be obtained in all (see Figure 3).

Insert Figure 3 here

Figure 3: Combinations of preferences

To elucidate the issue, we will begin by formulating the utility functions of the agents and calculating the optimal incentive system of the principal in the extreme cases of two-sided rivalry (*RR*) as well as two-sided altruism (*AA*). We will then proceed to examine and interpret in greater detail the impact of different types of social preferences exhibited by agents on the profit share rates, as established by the principal, as well as on the profitability of the firm in various scenarios. In our analysis, we will pay particular attention to the distinction between one-sided (*RE* or *ER* and *AE* or *EA*) and two-sided homogeneous (*RR* and *AA*) social preferences. These combinations will be considered separately, taking into account that the strength of both agents' social preferences can vary. After that, we will look into the case of two-sided heterogeneous social preferences (*RA* or *AR*). This approach enables us to draw fresh conclusions about the optimal composition of a team and the influence that large differences between the preferences of its members can have on firm profitability.

Prerequisites and Structure of the Basic Analytical Model

To analyze the impact of diverse preferences, we assume a LEN model with one principal P and two agents i ($i = A, B$), each of whom leads his or her own decentralized division. The profits x_i of the two decentralized divisions, which depend on the agents' non-observable efforts a and b as well as on the error terms ε_i , which represent stochastic environmental influences, add up to the firm profits, denoted by x . Therefore, the profit functions based on (isolated) production functions are:

$$x_A = a + \varepsilon_A \quad ; \quad x_B = b + \varepsilon_B. \quad (1)$$

The error terms are normally distributed with a mean of zero, variance σ_i^2 and correlation coefficient ρ . In this paper no technological dependences between the two decentralized divisions are assumed. The agents' efforts cause non-monetary quadratic disutility V_i , given by

$$V_A = \frac{1}{2} c_A a^2 \quad ; \quad V_B = \frac{1}{2} c_B b^2, \quad (2)$$

where the individual coefficients c_i allow for unequal marginal costs. Both agents are offered linear contracts, where the total amount of wage compensation S_i comprises a fixed payment α_0 (β_0) as well as a proportional fraction in the profits of each of the two decentralized divisions, which are determined by share rates α_A , α_B and β_A , β_B :

$$S_A = \alpha_0 + \alpha_A x_A + \alpha_B x_B \quad ; \quad S_B = \beta_0 + \beta_A x_A + \beta_B x_B. \quad (3)$$

The principal optimizes his or her objective function with respect to the wage compensation coefficients in (3). Depending on the endogenously determined values of α_B and β_A three different compensation schemes may develop: (a) individual compensation for $\alpha_B = 0$ ($\beta_A = 0$), (b) relative performance evaluation for $\alpha_B < 0$

($\beta_A < 0$) and (c) team-based compensation for $\alpha_B > 0$ ($\beta_A > 0$). The principal and the agents have exponential utility functions. Since the principal is assumed to be risk-neutral, he or she maximizes the expected value of his or her residuum after wage payments:

$$U_p = (1 - \alpha_A - \beta_A)x_A + (1 - \alpha_B - \beta_B)x_B - \alpha_0 - \beta_0. \quad (4)$$

Both agents are presumed to be strictly risk-averse. The strength of their risk aversion is measured by the constant coefficients $r_A > 0$ and $r_B > 0$, where higher values of r_i imply a higher degree of risk aversion.

Three types of preferences are taken into account: **pure self-interest**, **rivalry** and **altruism**. Agents who are driven by pure self-interest consider only their own material needs when choosing effort levels a and b , whereas competitive and altruistic agents also consider the effects of their decisions on the financial welfare of another agent. In the following analysis, it is assumed that social preferences refer only to the agents' remuneration and do not include parameters that are harder to observe such as costs of effort. By means of S_i, S_j ; $i, j = A, B$ and $i \neq j$, we therefore define:

$$F_i(S_i(\cdot), S_j(\cdot)) = S_i(\cdot) - k_i(l_i S_j(\cdot) - S_i(\cdot)) \text{ for rivalry} \quad (5)$$

and

$$G_i(S_i(\cdot), S_j(\cdot)) = m_i S_i(\cdot) + n_i S_j(\cdot) \text{ for altruism.} \quad (6)$$

Since both types of social preferences can vary in strength, the above specifications represent a continuum of different behavioral types, each of which includes **pure self-interest** as a special case for either $k_i = 0$ or $n_i = 0$ and $m_i = 1$:

$$H_i(S_i(\cdot), S_j(\cdot)) = S_i(\cdot). \quad (7)$$

In the case of **rivalry**, the social preference term $k_i(l_i S_j(\cdot) - S_i(\cdot))$ in equation (5) for agent i compares the other agent j 's realized remuneration with agent i 's own reward. If the resulting value is positive, agent i feels disadvantaged and suffers a disutility (envy). Therefore, the principal has to heighten agent i 's wage payments if he or she wants that agent to cooperate. On the other hand, if the result has a negative value, agent i perceives the other agent as being at a disadvantage, and draws additional utility from such a situation. This may be interpreted as "schadenfreude", so to speak, but also as pride in personal achievements. Consequently agent i accepts lower remuneration, which suggests that he or she works harder for unchanged incentives. The social preference parameter k_i ($k_i \geq 0$) in formula (5) measures the strength of an agent's rivalry and can take different values for each agent. Higher values indicate a stronger rivalry. The parameter l_i measures agent i 's aspiration level. It defines the ratio of the agents' remuneration S_i/S_j for which the effect of his or her social preference changes from utility-reducing to utility-enhancing. For $l_i = 1$, agent i is envious and suffers a disutility if he or she earns less than j . If the rewards of agent i exceed those of j , that agent feels more satisfied after evaluating his or her remuneration with relation to that of his or her peer. This evaluation regards utility gains, which are compared to the utility that agent would receive in case of purely selfish behavior. For values of $l_i > 1$, there exist allocations of rewards where agent i is not satisfied with his or her own remuneration although agent i earns more than j . In this case, this indicates that agent i does not accept that agent j should receive even a smaller reward (one might say, that agent i begrudges agent j 's reward). However, values of $l_i < 1$ indicate that agent i may be gleeful or proud evaluating his rewards relative to those

of the other agent although he or she earns less than j , which appears implausible. Therefore, we restrict the co-domain of l_i to the interval $l_i \in [1; \infty]$.

Unlike rivalry, which is considered above, **altruism**, as described by the formula given in equation (6), suggests that an agent is content when the other agent receives a greater monetary reward. This results in a willingness to exert more effort even when the first agent's own remuneration remains unchanged. The utility function is strictly increasing in both agents' wage compensations S_i and S_j . The weighting factor $0 \leq m_i \leq 1$ indicates that when an agent attaches greater weight to the other agent's remuneration through $n_i \geq 0$ the subjective importance of the first agent's own reward can decrease. Therefore, in equation (6) both agents strive to maximize the weighted sum of monetary payoffs. For each agent, the wage payments S_i and S_j are perfect substitutes.²⁵ Using the general specification (6), we can differentiate special cases of altruism for

$$m_i = m_j = 1 \quad (8a)$$

as well as

$$m_i + n_i = 1 \quad (8b)$$

For the latter, the social preference term in (6) can be rewritten as:

$$G_i(\cdot) = (1 - n_i)S_i(\cdot) + n_i S_j(\cdot) = S_i(\cdot) + n_i(S_j(\cdot) - S_i(\cdot)).^{26} \quad (9)$$

Edgeworth (1881) calls the relation of weighting factors $\frac{n_i}{m_i}$ and $\frac{n_i}{1 - n_i}$ "coefficients of

effective sympathy." They can be interpreted as measures for the strength of

altruism. When $\frac{n_i}{m_i} < 1$ ($\frac{n_i}{1 - n_i} < 1$), one agent wishes that the other agent receives

high remuneration, but still values his or her own rewards more. For $\frac{n_i}{m_i} > 1$

($\frac{n_i}{1-n_i} > 1$), however, the first agent is even willing to forgo part of his or her own remuneration if the other agent receives more instead. This form of altruism can also be called selfless behavior. The main difference between specifications (8a) and (8b) lies in a shift of the reference point. In the case of (8a), selfless behavior begins at $n_i = m_i = 1$, whereas in the case of (8b), it begins at $n_i = 0.5$.

The formulas that express different types of social preferences given in equations (5) and (6) are identical if the mathematical relations

$$1 + k_i = m_i \quad ; \quad -k_i l_i = n_i \quad (10)$$

hold. This formal equivalence allows us to compare utility values and derive certain implications concerning the optimal composition of teams. For the sake of simplicity, here we concentrate on clarifying the procedure that yields an optimal solution, mainly for cases of rivalry. Using equation (10), one can deduce the corresponding results for altruism and heterogeneous social preferences without further calculation.

4. Structure of the Optimal Incentive System

Given the stated assumptions, the agents' utility functions can be written using certainty equivalent notation. In the case of rivalry, one finds that:

$$CE_i = E[(1 + k_i) \cdot S_i(\cdot) - k_i l_i S_j(\cdot)] - V_i(\cdot) - \frac{r_i}{2} \text{Var}[(1 + k_i) \cdot S_i(\cdot) - k_i l_i S_j(\cdot)]. \quad (11)$$

Similarly, in the case of altruism we find that:

$$CE_i = E[m_i S_i(\cdot) + n_i S_j(\cdot)] - V_i(\cdot) - \frac{r_i}{2} \text{Var}[m_i S_i(\cdot) + n_i S_j(\cdot)]. \quad (12)$$

In either situation, both agents choose effort levels a and b in order to maximize their certainty equivalents (11) or (12), which, using equations (1), (2) and (3), can also be written as (the example refers to agent A in cases of rivalry):

$$\begin{aligned} \max_a CE_A = & (1+k_A)\alpha_0 + (1+k_A)\alpha_A a + (1+k_A)\alpha_B b - k_A l_A \beta_0 - \\ & k_A l_A \beta_A a - k_A l_A \beta_B b - \frac{1}{2} c_A a^2 - \frac{r_A}{2} \cdot \{ [\alpha_A (1+k_A) - k_A l_A \beta_A]^2 \sigma_A^2 + \\ & [\alpha_B (1+k_A) - k_A l_A \beta_B]^2 \sigma_B^2 + 2 \cdot [\alpha_A (1+k_A) - k_A l_A \beta_A] \cdot [\alpha_B (1+k_A) - k_A l_A \beta_B] \sigma_A \sigma_B \rho \}. \end{aligned} \quad (13)$$

The principal anticipates this behavior and restricts the optimization problem through the so-called incentive compatibility constraints. The corresponding reaction functions of the agents in the case of two-sided rivalry (*RR*) are:

$$a = \max \left[\frac{(1+k_A)\alpha_A - k_A l_A \beta_A}{c_A}; 0 \right]; \quad b = \max \left[\frac{(1+k_B)\beta_B - k_B l_B \alpha_B}{c_B}; 0 \right] \quad (14)$$

while for two-sided altruism (*AA*), they are:

$$a = \max \left[\frac{m_A \alpha_A + n_A \beta_A}{c_A}; 0 \right]; \quad b = \max \left[\frac{m_B \beta_B + n_B \alpha_B}{c_B}; 0 \right]. \quad (15)$$

The principal further has to consider the participation constraints, which suggest that each agent has to receive at least his or her reservation utility, normalized to zero in our model. Given these limitations, the principal maximizes his or her utility function (4), which for two-sided rivalry can also be written as

$$\begin{aligned} U_P = & (1 - \alpha_A - \beta_A) \left[\frac{(1+k_A)\alpha_A - k_A l_A \beta_A}{c_A} \right] + \\ & (1 - \alpha_B - \beta_B) \left[\frac{(1+k_B)\beta_B - k_B l_B \alpha_B}{c_B} \right] - \alpha_0 - \beta_0, \end{aligned} \quad (16)$$

over the wage compensation parameters $\alpha_0, \alpha_A, \alpha_B, \beta_0, \beta_A, \beta_B$. Performing the typical optimization steps yields:²⁷

$$\alpha_A = \frac{1+k_B}{1+k_B+k_B l_B} \cdot P_1 + \frac{k_A l_A}{1+k_A+k_A l_A} \cdot Q_1 \quad (17)$$

$$\alpha_B = \frac{1+k_B}{1+k_B+k_B l_B} \cdot P_2 + \frac{k_A l_A}{1+k_A+k_A l_A} \cdot Q_2 \quad (18)$$

$$\beta_A = \frac{1+k_A}{1+k_A+k_A l_A} \cdot Q_1 + \frac{k_B l_B}{1+k_B+k_B l_B} \cdot P_1 \quad (19)$$

$$\beta_B = \frac{1+k_A}{1+k_A+k_A l_A} \cdot Q_2 + \frac{k_B l_B}{1+k_B+k_B l_B} \cdot P_2 \quad (20)$$

with

$$P_1 = \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \quad ; \quad P_2 = -\frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \cdot \frac{\sigma_A}{\sigma_B} \rho \quad (21)$$

$$Q_2 = \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} \quad ; \quad Q_1 = -\frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} \cdot \frac{\sigma_B}{\sigma_A} \rho \quad (22)$$

The corresponding results for altruism and heterogeneous social preferences are obtained by applying the formal relations given in equation (10). From the reference case of purely selfish behavior, we know that an agent's wages are lower for higher values of his or her counterpart's performance measure. This so-called "relative performance evaluation" is useful in the presence of stochastic dependences because it enables us to filter out the risks that can be traced back to common occurrences that lie outside the agents' responsibilities and affect their remuneration equally (Holmström 1982). When social preferences are introduced, it can be seen from equations (17) through to (20) that the principal reacts to the agents' comparisons of remuneration by reallocating their variable wage compensation components through shifting the shares in performance measures that each agent receives, where α_A and β_A as well as α_B and β_B are interdependent. The weightings of performance measures are differentially divided between the two agents according to the type(s) of social preferences, which can cause major changes in the optimal design of the incentive system.²⁸ In order to clarify that argument, we shall analyze separately the cases of one- and two-sided social preferences.

4.1. Profit Share Rates in Cases of One-Sided Rivalry

Impact of Rivalry on the Agents' Shares in A's Profits

For the analysis of one-sided rivalry we assume that agent A behaves competitively while agent B is completely selfish ($RE : k_A > 0; k_B = 0$). The optimal expressions for share rates α_A and β_A in equations (17) and (19) can be reduced to:

$$\alpha_A = \frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A} + \frac{k_A l_A}{1 + k_A + k_A l_A} \cdot Q_1 \quad (23)$$

$$\beta_A = \frac{1 + k_A}{1 + k_A + k_A l_A} \cdot Q_1, \quad (24)$$

where $Q_1 < 0$. Thus, β_A in equation (24) simultaneously fulfills two complementary functions. It serves

- as an insurance parameter that reduces the wage compensation risk for agent B, and
- as a source of intrinsic motivation for competitive agent A.

Considering agent A's reaction function in equation (14), one can observe that *negative values* of β_A have a performance-enhancing effect for that agent. This is due to a reduction in agent B's wage compensation, which leads to greater rivalry and thus higher motivation for A. The increased effort that A makes as a result leads to an increase in the cost of effort. Thus, the principal must reduce A's direct incentive intensity by reducing α_A . In turn, this also reduces the relevant wage compensation risk for agent A. Accordingly, equation (23) can be written as:

$$\alpha_A = \frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A} + \frac{k_A l_A}{1 + k_A} \cdot \beta_A, \quad \beta_A < 0. \quad (25)$$

At the same time, because $\beta_A < 0$ as an insurance parameter for agent B no longer has only a risk-reducing effect when taking rivalry into account (RE or RR), the

absolute value of β_A becomes lower, compared to the reference point of purely selfish behavior:

$$Risk_A = \frac{r_A}{2} \cdot \left\{ [\alpha_A(1+k_A) - \mathbf{k}_A \mathbf{l}_A \beta_A]^2 \sigma_A^2 + [\alpha_B(1+k_A) - k_A l_A \beta_B]^2 \sigma_B^2 + 2 \cdot [\alpha_A(1+k_A) - \mathbf{k}_A \mathbf{l}_A \beta_A] \cdot [\alpha_B(1+k_A) - k_A l_A \beta_B] \sigma_A \sigma_B \rho \right\}. \quad (26)$$

When we consider the terms in bold in equation (26),²⁹ it becomes obvious that $\beta_A < 0$ on the one side may reduce the risk carried by agent B, but at the same time has a risk-increasing effect for agent A. The latter is more pronounced for increasingly negative β_A values. Less negative values of β_A reduce the risk for agent A. Consequently, when determining β_A , it is necessary for the principal to trade off between optimal incentive intensity on one side and optimal risk-sharing on the other. The difference from the reference case of completely selfish agents is that, in this case (*RE*), this necessity arises only indirectly as a result of agent A's social preferences. As a consequence, the absolute value of negative β_A shrinks as the values of the social preference parameter k_A , as well as of the aspiration level l_A , rise, while the positive values of α_A are simultaneously reduced.

Impact of Rivalry on the Agents' Shares in B's Profits

Again, assuming *RE* : $k_A > 0; k_B = 0$, the optimal expressions for share rates α_B and β_B in equations (18) and (20) become:

$$\alpha_B = -\frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \cdot \frac{\sigma_A}{\sigma_B} \rho + \frac{k_A l_A}{1+k_A+k_A l_A} \cdot Q_2 \quad (27)$$

$$\beta_B = \frac{1+k_A}{1+k_A+k_A l_A} \cdot Q_2. \quad (28)$$

Unlike in the previous analysis concerning β_A , here when the principal determines α_B he or she needs to trade off two opposite effects:

- filtering out risks in the incentive system of agent A, and
- reducing the possible negative consequences caused by agent A's rivalry.

From previous agency models, it is known that stochastic dependences further the use of *relative performance evaluation*, which yields negative values for α_B , while agent A's rivalry has the opposite effect by enhancing the application of *team-based compensation*. Thus, when both forms of interdependences occur at the same time it should be determined which of these two basic types of compensation will be favored. Figure 4 illustrates the range of values for share rate α_B as a function of the strength of A's rivalry k_A ³⁰ on one side, and the strength of the stochastic dependence ρ on the other.³¹ The graph reveals that the horizontally drawn zero plane and the curved α_B surface intersect for every value of A's rivalry k_A . Consequently, for every strength of A's rivalry k_A , there is a threshold value $\rho^*(k_A)$ where agent A's share rate α_B switches signs. The intersecting line $\rho^*(k_A)$ is strictly increasing and concave, meaning that the correlation's threshold value $\rho^*(k_A)$ rises as the influence of rivalry k_A increases. The analysis makes clear that higher values of ρ augment the merits of relative performance evaluation, while team-based pay becomes more favorable when agent A's rivalry is more strongly developed. If $\rho < \rho^*(k_A)$, the random variation of both agents' performance measures is affected only slightly by common events.

Insert Figure 4 here

Figure 4: Range of values for α_B , dependent on the correlation coefficient ρ as well as the strength of rivalry k_A .

Thus, relative performance evaluation is not a very powerful means of reducing the risk carried by agent A. In addition to the comparatively small risk-reduction potential for $\alpha_B < 0$, increasing α_B above zero enables the principal to ensure that agent A will always benefit when higher wage payments are granted to agent B.³² This measure,

which arises from A's rivalry, increases this agent's utility and therefore leads to higher firm profits. As a consequence, the principal implements a team-based wage compensation scheme. However, if the threshold value of the correlation coefficient $\rho^*(k_A)$ is exceeded, the positive risk-reducing effects are greater than the potentially negative consequences of rivalry, and thus relative performance evaluation becomes advantageous.

As a last result of agent A's rivalry, and of the consequential reallocations of the agents' variable wage compensation components as incentives, the principal lowers *incentive intensity* β_B of agent B. Accordingly, the sum of β_B in equation (28) and the second summand of the mathematical expression for α_B in equation (27) are identical to agent B's incentive intensity in the case of purely selfish behavior.

4.2. Profit Share Rates in the Case of One-Sided Altruism

Impact of Altruism on the Agents' Shares in A's Profits

The analysis in cases of one-sided altruism, which is characterized by one altruistic and one selfish agent (*AE*), suggests that there are many analogies between one-sided altruism (*AE*) and one-sided rivalry (*RE*). In the following, we shall point out the major differences between these two cases (*AE* and *RE*) that arise in the interpretation of our results. Here too, share rate β_A simultaneously fulfills the same two functions as in the situation of rivalry (i.e. as an insurance parameter that reduces the wage compensation risk for agent B, and as a source of intrinsic motivation for the competitive agent A). However, in the case of one-sided altruism, the effects of β_A work in the opposite direction: negative values of β_A reduce the risk for agent B, but at the same time, these values lower the intrinsic motivation of agent A.³³ To see why this is the case, we shall turn to the definition of altruism introduced further up: this type of social preference characteristically includes the willingness to

bear a personal cost for actions that increase the prosperity of another person. Therefore, positive values of β_A would translate into greater effort from altruistic agent A, as this is expected to lead both to higher rewards for agent A and better remuneration for agent B.

The opposite applies when β_A takes negative values, as in our case, because in that case an increase in the effort of agent A will improve A's financial rewards but reduce agent B's expected income. Thus, altruism partly undermines the positive outcome of hard work for agent A, since, as a consequence of this social preference, that agent internalizes the negative external effect of his or her actions on B and therefore also suffers from their negative impact on B's remuneration. This results in less willingness to bear the cost of personal effort, than in the situation of purely selfish behavior. As a countermeasure, the principal has to boost agent A's explicit incentives by increasing α_A , otherwise A's effort will fall short of the principal's expectations. These observations are more noticeable at higher values of the "coefficient of effective sympathy" $\frac{n_A}{m_A}$.

Impact of Altruism on the Agents' Shares in B's Profits

Agent A's altruism leads to an increase in explicit incentives for agent B by increasing his or her variable wage compensation component β_B . This is because the altruistic agent's utility is positively affected by the other agent's total remuneration, as a result of which the altruistic agent is willing to forgo his or her own financial reward. Thus, the altruistic agent's variable wage payment can be reduced by decreasing α_B . Consequently, the mathematical relation

$$\alpha_B = -\frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A} \cdot \frac{\sigma_A}{\sigma_B} \rho - \frac{n_A}{m_A} \cdot \beta_B \quad , \quad \beta_B > 0 \quad (29)$$

holds. The principal's scheme can be configured to subsidize the greater effort of the

selfish agent by exploiting the other agent's altruism. The resulting higher risk premiums are compensated through the fixed remuneration components. Reducing m_A from 1 and raising n_A from 0, i.e., increasing the “coefficient of effective sympathy” $\frac{n_A}{m_A}$, amplifies the described effects, as can be seen from equation (29).

In contrast to the rivalry situation, here α_B is negative, even in the absence of social preferences. Therefore, there is no need for a trade-off between relative performance evaluation and team-based compensation. Both agents are rewarded relatively to one another, irrespective of the strength of agent A's altruism and the degree to which the error terms in their performance measures are correlated.

4.3. Combined Effects of Two-Sided Social Preferences

If both agents take each other's wage payments into account, their respective social preferences simultaneously affect all four share rates, which results in overlapping effects. Table 1 shows that a more developed sense of rivalry on the part of both agents (RR) operates in the same direction for all four share rates.

	Agent A		Agent B	
	α_A	α_B	β_A	β_B
FOC	$\frac{\partial \alpha_A}{\partial k_A}, \frac{\partial \alpha_A}{\partial k_B} < 0$	$\frac{\partial \alpha_B}{\partial k_A}, \frac{\partial \alpha_B}{\partial k_B} > 0$	$\frac{\partial \beta_A}{\partial k_A}, \frac{\partial \beta_A}{\partial k_B} > 0$	$\frac{\partial \beta_B}{\partial k_A}, \frac{\partial \beta_B}{\partial k_B} < 0$
$k_A \uparrow$	↓	↑	↑	↓
$k_B \uparrow$	↓	↑	↑	↓

Table 1: Impact on share rates caused by a change in weighting for each agent's rivalry

This suggests that the agents' direct incentive intensities α_A and β_B are reduced when the social preferences of both are more developed. At the same time, each

agent's shares in the other agent's performance measure grow with an increase in the weighting of rivalry, which is represented by $k_A \uparrow$ and $k_B \uparrow$. Thus, the advantages of team-based compensation as compared to relative performance evaluation are further enhanced in a situation of two-sided rivalry. The impact of varying the intensity of both agents' altruism (AA) on the four share rates is depicted in Table 2.

	Agent A		Agent B	
	α_A	α_B	β_A	β_B
FOC	$\frac{\partial \alpha_A}{\partial m_A}, \frac{\partial \alpha_A}{\partial m_B} < 0$ $\frac{\partial \alpha_A}{\partial n_A}, \frac{\partial \alpha_A}{\partial n_B} > 0$	$\frac{\partial \alpha_B}{\partial m_A}, \frac{\partial \alpha_B}{\partial m_B} > 0$ $\frac{\partial \alpha_B}{\partial n_A}, \frac{\partial \alpha_B}{\partial n_B} < 0$	$\frac{\partial \beta_A}{\partial m_A}, \frac{\partial \beta_A}{\partial m_B} > 0$ $\frac{\partial \beta_A}{\partial n_A}, \frac{\partial \beta_A}{\partial n_B} < 0$	$\frac{\partial \beta_B}{\partial m_A}, \frac{\partial \beta_B}{\partial m_B} < 0$ $\frac{\partial \beta_B}{\partial n_A}, \frac{\partial \beta_B}{\partial n_B} > 0$
$n_A/m_A \uparrow$	$\uparrow$	$\downarrow$	$\downarrow$	$\uparrow$
$n_B/m_B \uparrow$	$\uparrow$	$\downarrow$	$\downarrow$	$\uparrow$

Table 2: Impact on share rates caused by a change in weighting for each agent's altruism

It becomes apparent that both agents' "coefficients of effective sympathy", which are used as a measure of the strength of the altruism each agent exhibits, also operate in the same direction for all four share rates. Accordingly, the stronger each agent's altruism becomes, the more that agent's own performance measure counts in his or her own wage compensation system. This is because with increasing altruism, an agent attaches greater weight to the other agent's remuneration than to his or her own personal reward in relative terms. By contrast, each agent's share in the other's performance measure is increasingly reduced as his or her "coefficient of effective sympathy" adopts higher values. Finally, the impact of varying heterogeneous social preferences, i.e., altruism for agent A and rivalry for agent B (AR), is shown in Table 3.

	Agent A		Agent B	
	α_A	α_B	β_A	β_B
FOC	$\frac{\partial \alpha_A}{\partial m_A}, \frac{\partial \alpha_A}{\partial k_B} < 0$ $\frac{\partial \alpha_A}{\partial n_A} > 0$	$\frac{\partial \alpha_B}{\partial m_A}, \frac{\partial \alpha_B}{\partial k_B} > 0$ $\frac{\partial \alpha_B}{\partial n_A} < 0$	$\frac{\partial \beta_A}{\partial m_A}, \frac{\partial \beta_A}{\partial k_B} > 0$ $\frac{\partial \beta_A}{\partial n_A} < 0$	$\frac{\partial \beta_B}{\partial m_A}, \frac{\partial \beta_B}{\partial k_B} < 0$ $\frac{\partial \beta_B}{\partial n_A} > 0$
$n_A/m_A \uparrow$	$\uparrow$	$\downarrow$	$\downarrow$	$\uparrow$
$k_B \uparrow$	$\downarrow$	$\uparrow$	$\uparrow$	$\downarrow$

Table 3: Impact on share rates caused by a change in weighting with regard to the agents' heterogeneous social preferences

In this case, both agents' social preferences influence the values of all four share rates in opposite directions. Regardless of this, the altruistic agent A obtains a contract based on a relative performance evaluation. For the competitive agent B, however, again, it is necessary to trade off relative performance evaluation against team-based compensation. As in the cases of one- and two-sided rivalry discussed previously (*RE* and *RR*), here too β_A is positive (and consequently team-based compensation should be used) if the correlation coefficient falls below a threshold

value $\rho^*\left(\frac{n_A}{m_A}, k_B\right)$. Conversely, for $\rho > \rho^*\left(\frac{n_A}{m_A}, k_B\right)$, relative performance evaluation

becomes optimal. The threshold value $\rho^*\left(\frac{n_A}{m_A}, k_B\right)$ increases with the strength of

agent B's rivalry and for decreasing values of the "coefficient of effective sympathy"

$\frac{n_A}{m_A}$ exhibited by agent A. Compared to the situation described in section 4.1, here

the relation $\rho^*\left(\frac{n_A}{m_A}, k_B\right) < \rho^*(k_B)$ holds true, meaning that for the competitive agent

B, the benefits of relative performance evaluation increase in the case of heterogeneous social preferences, as compared to a situation of one-sided rivalry.

5. The Influence of Social Preferences on the Profitability of the Firm

This section examines the maximum attainable firm profits for all possible combinations of the different types of (social) preferences considered in our model. Our aim is to shed light on two interrelated issues. First, we will identify the conditions under which the principal can make use of the agents' social preferences in each situation. Then, building on the results, we will turn to the central concern of our paper and broach the issue of the optimal combination of different types of (social) preferences for which firm profits are maximized. In this analysis, we assume that the principal optimally designs the wage compensation system according to the principles outlined above. Therefore, the maximum value of his or her objective function is calculated by assigning to the share rates in equation (16) their optimal values given in equations (17) to (22). This yields:³⁴

$$U_P^*(k_A, k_B) = \frac{1}{2c_A} \cdot \frac{(1+k_A)(1+k_B) - k_A l_A k_B l_B}{1+k_B + k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} + \frac{1}{2c_B} \cdot \frac{(1+k_A)(1+k_B) - k_A l_A k_B l_B}{1+k_A + k_A l_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B}. \quad (30)$$

For the sake of clarity, we first consider the case of one-sided social preferences *RE* and *AE* in sections 5.1 and 5.2. We then go on to examine the optimal combinations in the case of two-sided homogeneous and heterogeneous social preferences in sections 5.3 (*RR* as well as *AA*) and 5.4 (*AR*).

5.1. The Profitability of One-Sided Rivalry

In the case of one-sided rivalry (*RE* : $k_A > 0; k_B = 0$) the optimal value of the objective function (30) reduces to:

$$U_P^*(k_A) = \frac{1}{2c_A} \cdot (1+k_A) \cdot \frac{1}{1+r_A\sigma_A^2(1-\rho^2)c_A} + \frac{1}{2c_B} \cdot \frac{1+k_A}{1+k_A+k_A l_A} \cdot \frac{1}{1+r_B\sigma_B^2(1-\rho^2)c_B}. \quad (31)$$

The first term in (30) and (31) expresses the principal's utility for agent A; the second term, the utility for agent B. Thus, agent A's rivalry influences the principal's economic prosperity in two ways. The first term is raised by the factor $(1+k_A) > 1$, since agent A, because of his or her rivalry, is intrinsically motivated to distinguish himself or herself from agent B in terms of wage payments, which leads agent A to make a greater effort. For that reason, we call this effect, which is represented by the factor $(1+k_A)$, the *intrinsic motivation effect*. This has desirable consequences on the trade-off between efficient incentives and efficient risk-sharing from the firm's point of view. However, the second term is reduced by the multiplier $\frac{1+k_A}{1+k_A+k_A l_A} < 1$. As is shown in section 4.1, the principal needs to reduce agent B's participation in his or her own performance measure β_B in order to react optimally to agent A's rivalry. The resulting lower incentive intensity for agent B causes that agent's extrinsic motivation to decrease and therefore his or her effort to fall below the level identified in the reference case of purely selfish behavior. The principal's welfare is therefore negatively affected by this so-called *extrinsic motivation effect*, which becomes more pronounced as the strength of rivalry k_A and the aspiration level l_A take higher values.

The total impact of agent A's rivalry is determined by the question of which of the two effects described above dominates. However, to answer this question we have to consider not only the relative impact of the factors influenced by social preferences, but also the contributions that both agents would make in a situation of completely selfish behavior. Therefore, we differentiate between two cases in the

following analysis: in the first case, we assume that the agents are identical except for their basic types of (social) preferences. Figure 5 illustrates the optimal value of the principal's objective function, subject to the strength of agent A's rivalry as well as the strength of the stochastic dependence.³⁵

Insert Figure 5 here

Figure 5: Impact of one-sided rivalry on firm profits when both agents are identical apart from their types of (social) preferences

It becomes clear, then, that the principal benefits from agent A's rivalry when the two agents are, apart from their social preferences, very similar in terms of their effort costs as well as their attitudes and exposure to risk.³⁶ However, as indicated by Figure 6, which represents the second case,³⁷ this need not be true if the self-centered agent B is significantly more important to the principal in a situation of purely selfish behavior (*EE*). This could occur if, for example, agent A is much more risk averse ($r_A \gg r_B$), has a more volatile performance measure ($\sigma_A \gg \sigma_B$), or suffers higher costs from comparable efforts ($c_A \gg c_B$) than agent B. The correlation ρ has no impact on the principal's opportunity to profit from agent A's rivalry (everything else being equal), since it influences equally the utility that can be extracted from both agents. All other parameters, exhibit threshold values which signify changes in the principal's opportunity of benefiting from agent A's rivalry.³⁸

The intuition for these formally derived results is as follows; If competitive agent A contributes to firm profits as much as or more than selfish agent B, then agent A's ambition amplifies his or her effectiveness and so the corresponding reduction in incentive intensity for B is of comparably minor importance.

Insert Figure 6 here

Figure 6: Impact of one-sided rivalry on firm profits when greater importance is attached to selfish agent B than to competitive agent A

By contrast, if agent B is of greater importance to the superordinate's goal, A's ambition can prove detrimental to firm profit as, in that case, the principal is expected to privilege agent A over B in terms of shares in performance measure x_B , which would decrease the effort that the more capable agent B makes. Therefore, competitive behavior on the part of agents who make only small contributions to firm profits can have a negative effect, as this behavior requires the principal to pay those agents greater attention and discriminate against other high performers. The previous analysis leads to the following proposition:

Proposition 1 (Firm Profitability in the Case of One-Sided Rivalry): *The principal profits from an agent's rivalry if he or she chooses variable wage compensation components according to equations (17)–(22) (choice being an endogenous parameter) and if the exogenous conditions*

- a) *that the selfish agent's contribution to his or her utility does not exceed by far the contribution of the competitive agent and*
 - b) *that the competitive agent's strength of rivalry does not fall below a critical value*
- are fulfilled. In this case, the principal profits more from stronger rivalry.³⁹ One-sided rivalry necessarily proves detrimental to the principal if both exogenous conditions are simultaneously not fulfilled.*

Finally, Figure 5 also illustrates that the interrelationship between firm profits and the strength of agent A's rivalry is asymptotically linear, beginning with small values of k_A . As a consequence, its marginal impact is approximately constant. Furthermore, the slope of all curves $U_P^*(k_A)$ grows with increasing values for the correlation coefficient ρ , meaning that the principal profits more from agent A's rivalry when the stochastic dependence is stronger. This is caused by higher-order effects, which

arise from the multiplicative composition of both profit-enhancing effects (risk reduction through higher ρ and benefits caused by A's rivalry) in the principal's objective function. From an economic point of view, this result can be traced back to the complementarity of the functions described in section 4.

5.2. *The Profitability of One-Sided Altruism*

In general, altruism is usually perceived as a positive trait, since it involves concern for the prosperity of other people without resentment for their success. Accordingly, the altruistic agent enhances his or her efforts if the other agent's share in his or her own performance measure becomes larger, as can be seen from that agent's reaction functions in equation (15). The altruistic agent is therefore willing to bear additional effort costs if the other agent's expected remuneration rises. Bearing that in mind, this section examines whether the principal can always profit from an agent's altruistic behavior or if there is a trade-off similar to that observed in the analysis of rivalry described above. In the case of altruism, the optimal value of the principal's objective function, which is identical to firm profits, becomes:

$$U_P^*(m_A, n_A, m_B, n_B) = \frac{1}{2c_A} \cdot \frac{m_A m_B - n_A n_B}{m_B - n_B} \cdot \frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A} + \frac{1}{2c_B} \cdot \frac{m_A m_B - n_A n_B}{m_A - n_A} \cdot \frac{1}{1 + r_B \sigma_B^2 (1 - \rho^2) c_B}. \quad (32)$$

For one-sided altruism on the part of agent A ($AE : m_B = 1, n_B = 0$), this expression reduces to:

$$U_P^*(m_A, n_A) = \frac{1}{2c_A} \cdot m_A \cdot \frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A} + \frac{1}{2c_B} \cdot \frac{m_A}{m_A - n_A} \cdot \frac{1}{1 + r_B \sigma_B^2 (1 - \rho^2) c_B}. \quad (33)$$

Again, two distinct effects that affect the principal's profit must be noted: the additional weighting factor $0 < m_A < 1$ in agent A's contribution describes the *intrinsic*

motivation effect, which, in the case of altruism, leads to a reduction of the principal's profits. This effect occurs only if agent A, as a result of altruism, attaches less importance to his or her own wage payments than he or she would if self-interest dictated his or her behavior. Increasing α_A then adds fewer incentives for agent A to exhibit greater effort than in the situation of completely selfish behavior. This makes it increasingly difficult for the principal to motivate altruistic agent A by means of variable wage components. Consequently, the less agent A cares about his or her own remuneration, the more that agent's contribution to firm profits decreases. However, it remains unaltered if the weight that agent A, despite his or her altruism, attaches to his or her own monetary compensation stays constant.

The weighting factor $\frac{m_A}{m_A - n_A} > 1$ in agent B's contribution to firm profits determines the strength of the *extrinsic motivation effect*. As stated in section 4.2 the principal makes use of agent A's altruism to enhance agent B's effort by raising the latter's incentive intensity β_B . As a result, the more similar the values that agent A attaches to his or her own wage payments and those of agent B, the greater the increase in firm profits. Taking $n_A \rightarrow m_A$, one obtains

$$\lim_{n_A \rightarrow m_A} U_P^*(m_A, n_A) = \lim_{n_A \rightarrow m_A} \left[\frac{1}{2c_A} \cdot m_A \cdot \frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A} + \frac{1}{2c_B} \cdot \frac{m_A}{m_A - n_A} \cdot \frac{1}{1 + r_B \sigma_B^2 (1 - \rho^2) c_B} \right] = +\infty \quad (34)$$

which indicates that in the extreme case when the altruistic agent attaches nearly the same value to both agents' profits in his or her own utility function, this could lead to a theoretically unlimited increase in firm profits.

We continue with a two-step analysis of the total effect that one-sided altruism has on firm profits: first, we observe that in situations when agent A, despite his or

her altruism, attaches constant weight to his or her own wage payments ($m_A = 1$), it is clear that there are no effects operating in opposite directions and the principal can hence always benefit from that agent's social preference. Thus, the principal's gains increase in proportion to the value that the altruistic agent attaches to the other agent's prosperity with regards to his or her own remuneration. Second, we see that negative consequences on the achievement of objectives are likely to occur only when altruistic agent A contributes much more to firm profits than selfish agent B. A precondition for this, however, is that agent A attaches much less weight to his or her own wage payments than he or she would in the case of completely selfish behavior ($m_A < 1$), and, at the same time, that the relative importance of both agents' compensation does not become too similar ($n_A < m_A$). On the basis of the above, we can reach the following proposition:

Proposition 2 (Firm Profitability in the Case of One-Sided Altruism): *In all cases the principal's profits grow if the altruistic agent attaches greater weight to the other (selfish) agent's wage payments. A negative impact on the principal's profits can only occur in the rare cases where the altruistic agent*

- a) contributes more than the selfish one,*
- b) simultaneously attaches much less weight to his or her own wage payments than he or she would in cases of pure self-interest and*
- c) at the same time does not attach the same, or almost the same weight to the remuneration of both agents in his or her utility function.*

The current model specifications do not allow us to make general statements about whether, from a firm perspective, it is easier to benefit from an agent's rivalry rather than altruism.

5.3. The Influence of Two-Sided Homogeneous (Social) Preferences on Firm Profits

In cases of one-sided social preferences it was shown that the principal can usually benefit from an agent's social behavior. The next two sections extend that analysis by assuming that both agents exhibit social preferences. First, we examine homogeneous social preferences RR or AA , then we turn to the case of heterogeneous social preferences AR . We proceed by breaking down the weighting factors in the principal's optimal goal value that are determined by social preferences (see equations [30] and [32]) in order to examine separately each of their components. We conclude by an analysis of the total effect.

First, consider agent A's weighting factor in the case of two-sided **rivalry**

$\frac{(1+k_A)(1+k_B)-k_A l_A k_B l_B}{1+k_B+k_B l_B}$ in equation (30). Its *first term* $\frac{(1+k_A)(1+k_B)}{1+k_B+k_B l_B}$ can be subdivided

into *intrinsic* and *extrinsic motivation effects* that operate in opposite directions. As is

clear from the analysis in section 4.1, the *extrinsic motivation effect* $\frac{1+k_B}{1+k_B+k_B l_B}$

stems from the reallocation of variable wage compensation components between the

two agents by the principal because of agent B's rivalry. As a consequence, agent B

is favored over agent A. The resulting weakened incentive intensity α_A for agent A

leads that agent to reduce his or her effort. However, at the same time, agent A's

own ambition to achieve a result that is superior to B's, which is formally represented

by the factor $(1+k_A)$, in turn promotes his or her effort. Therefore, in the case of two-

sided social preferences, the *extrinsic* and the *intrinsic motivation effect* have a

mutual impact on each other with regard to firm profits. The second term $\frac{-k_A l_A k_B l_B}{1+k_B+k_B l_B}$

consists of the mathematical components $-k_A l_A$ and $\frac{k_B l_B}{1+k_B+k_B l_B}$. Together they form

an additional constituent of the *intrinsic motivation effect*, which exists only in the situation of two-sided social preferences. The principal's reallocation of variable wage compensation components as a response to agent B's rivalry leads share rate β_A to increase by $\frac{k_B l_B}{1 + k_B + k_B l_B} \cdot \frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A}$. Thus, agent B is favored over agent A. In its turn, A's own rivalry prompts him or her to differentiate himself or herself as much as possible from agent B in terms of wage compensation. The preferential treatment of agent B consequently lowers agent A's intrinsic motivation, as can be seen from the subtrahend $\frac{-k_A l_A}{c_A} \beta_A$ in agent A's reaction function (see equation [14]). As a result, agent A reduces his or her efforts. This reduction increases not only in both agents' strength of rivalry k_A and k_B , but also in both agents' aspiration levels l_A and l_B , since higher values of l_B lead to a further increase in β_A , to which agent A attaches more significance for higher values of l_A .

Before examining the total effect, we shall analyze agent A's weighting factor in the situation of two-sided **altruism** in equation (32) $\frac{m_A m_B - n_A n_B}{m_B - n_B}$. Here, the minuend $\frac{m_A m_B}{m_B - n_B}$ again consists of an *intrinsic* as well as an *extrinsic motivation effect*. In this case, however, agent B's altruism leads to an increase in agent A's incentive intensity α_A , which is reflected in the factor $\frac{m_B}{m_B - n_B} > 1$. A's own altruism, however, leads to reduced intrinsic motivation as that agent attaches less importance to his or her own wage compensation ($m_A < 1$). The subtrahend $-\frac{n_A n_B}{m_B - n_B}$ comprises the two components n_A and $-\frac{n_B}{m_B - n_B}$. Because of agent B's altruism, that agent's

share rate β_A is reduced by $-\frac{n_B}{m_B - n_B} \cdot \frac{1}{1 + r_A \sigma_A^2 (1 - \rho^2) c_A}$, which causes the intrinsic motivation of agent A to decrease. This effect increases as the importance that agent A, due to his own altruism, places on agent B's wage compensation increases ($n_A \uparrow$).

Insert Figure 7 left and Figure 7 right side-by-side here

Figure 7: Impact of two-sided rivalry and two-sided altruism on firm profits

The same factors used in the analysis above can be identified in agent B's weighting factors, $\frac{(1 + k_A)(1 + k_B) - k_A l_A k_B l_B}{1 + k_A + k_A l_A}$ and $\frac{m_A m_B - n_A n_B}{m_A - n_A}$ respectively. To examine the total effect, in Figure 7 we depict firm profits as a function of the strength of both agents' homogeneous social preferences in cases of rivalry and in cases of altruism.⁴⁰ Figure 7 illustrates that equally developed homogeneous social preferences cannot be beneficial from a firm perspective. In either case, firm profits are maximized when only one agent has a predisposition, preferably strong, for either rivalry or altruism. The intuitive explanation for this result is that the principal can only make use of the agents' social preferences if he or she reacts adequately to their social behavior by reallocating variable wage compensation components between the two agents. However, the measures that a principal is required to take in response to each agent's social preferences are interdependent. For example, agent A's rivalry means that the principal increases that agent's share α_B in agent B's performance measure x_B at the expense of a reduced incentive intensity β_B for B. The stronger the social preferences of agent A, the more intense the principal's response. In this situation, if B behaves competitively as well, the principal's response to agent A's rivalry simultaneously has negative effects not only on agent B's extrinsic but also on his intrinsic motivation. This implies that the principal cannot achieve an advantageous

trade-off by means of an adequately designed wage compensation system in a case where both agents simultaneously behave competitively. Every reallocation of the variable wage compensation components leads to an improvement on one side but, at the same time, has a negative effect on the other.

The same argument holds true for a situation with two-sided altruism. In this case, the principal can only profit from each agent's altruism if he or she increases the other agent's incentive intensity while simultaneously reducing the altruist's share in the same performance measure. This reduction, however, decreases the other agent's intrinsic motivation if he or she also behaves altruistically. Therefore, the positive effects of the heightened incentive intensity are partly wasted. Again, the reallocation of the variable wage compensation components that has been caused by one agent's altruism has negative effects on the other agent if he or she exhibits the same degree of altruism, since neither agent wants to be privileged over the other. In both cases, the principal's actions in response to each of the agents' social preferences contradict one another, which means that the principal cannot benefit from the behavior that results from those actions if the social preferences of both agents are of equal strength.⁴¹ This analysis leads to the following proposition:

Proposition 3 (Firm Profitability in the Case of Two-Sided Homogeneous Social Preferences): *If agents are endowed with the same characteristics, firm profits are the same independently of the agents' types of preferences. From a firm's point of view, in such cases there is no advantage of social over selfish preferences. On the contrary, compared to the case of pure self-interest, here firm profits are higher if only one of the two agents exhibits a particular social preference (either altruism or rivalry), preferably to a high degree, while the other behaves in a purely selfish manner.*

However, higher aspiration levels l_A and l_B always lead to reduced effort on the part of both agents and therefore have a negative impact on firm profitability. This suggests that when agents have high expectations of their personal performances, this can actually have a negative effect on their output, because it may undermine employee satisfaction. Therefore, from a firm's point of view, agents with low aspiration levels are generally more productive and hence more beneficial for the firm. If the principal could choose which agent should exhibit social preferences, ideally it would be the agent who also exhibits smaller marginal costs, smaller risk aversion and a less variable performance measure. This agent would make a larger contribution to firm profits, even in the case of purely selfish behavior, and this effect would be further amplified by the principal's reallocation of variable wage compensation components as a response to the agents' social preferences.

5.4. The Profitability of Two-Sided Heterogeneous Social Preferences for the Firm

In the previous analysis, firm profits reached their maximum level when only one agent had a predisposition, preferably a strong one, for either rivalry or altruism, and the homogeneous social preference of the other agent was zero, meaning that the second agent was completely selfish. Here we examine whether firm profitability can be further enhanced if the two agents exhibit heterogeneous social preferences. The principal's optimal objective function value takes the form:

$$U_P^{**}(m_A, n_A, k_B) = \frac{1}{2c_A} \cdot \frac{m_A(1+k_B) + n_A k_B l_B}{1+k_B + k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} + \frac{1}{2c_B} \cdot \frac{m_A(1+k_B) + n_A k_B l_B}{m_A - n_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B}. \quad (35)$$

We begin by breaking down the various effects contained in the weighting factors

$$\frac{m_A(1+k_B) + n_A k_B l_B}{1+k_B + k_B l_B} \text{ (agent A) and } \frac{m_A(1+k_B) + n_A k_B l_B}{m_A - n_A} \text{ (agent B) then go on to consider}$$

their combined impact on firm profits. The weighting factor in the contribution of the **altruistic agent A** in (35) has two components. The *first term* $\frac{m_A(1+k_B)}{1+k_B+k_B l_B}$ can be broken down into two subcomponents, m_A and $\frac{1+k_B}{1+k_B+k_B l_B}$. As a result of agent B's rivalry, the principal shifts variable wage compensation away from A and toward B. Thus, A's incentive intensity α_A is reduced by the factor $\frac{1+k_B}{1+k_B+k_B l_B} < 1$. This negative *extrinsic motivation effect*, which lowers the effort of agent A, is reinforced by A's altruism if that agent attaches less weight to his or her own monetary incentives ($m_A < 1$). Overall, the first component of agent A's weighting factor has a negative impact on firm profits.

The *second term* $\frac{n_A k_B l_B}{1+k_B+k_B l_B}$ reflects an *intrinsic motivation effect*. Because of agent B's rivalry, β_A , which denotes how closely agent B's payments are linked to agent A's performance measure x_A , takes higher values than in the case of pure self-interest. This results in higher intrinsic motivation for the altruistic agent A. In other words, due to the increase in β_A , agent A's actions not only affect his or her own wage payments positively, but also increase those of agent B, to whom agent A attributes the improvement of his or her own financial situation as compared to the case of pure self-interest. Therefore, the more agent A takes agent B's wage payments into consideration ($n_A \uparrow$), the more willing he or she is to bear higher costs and make a greater effort. Thus, the second term affects firm profits positively.

The weighting factor in the contribution of the **competitive agent B** in (35) also has two components. The *first part* $\frac{m_A(1+k_B)}{m_A - n_A}$ is made of two subcomponents. The factor $(1+k_B)$ depicts an *intrinsic motivation effect* for agent B. Spurred by rivalry,

agent B wants to be as far ahead of agent A as possible. The principal exploits this ambitious behavior. Simultaneously, β_B is increased by the factor $\frac{m_A}{m_A - n_A} > 1$ because of agent A's altruism, which leads the principal to increase the direct incentive intensity for agent B. This further enhances agent B's (extrinsic) motivation.

The second part $\frac{n_A k_B l_B}{m_A - n_A}$ describes the *intrinsic motivation effect* for agent B.

Agent A's altruism ties his or her wage payments negatively to agent B's performance measure x_B ($\alpha_B < 0$). Therefore, A's rewards decrease as agent B enhances his or her efforts. This means that in such a situation competitive agent B draws additional utility from his or her own efforts, as can be seen from the mathematical expression $-\frac{k_B l_B}{c_B} \alpha_B$ in agent B's reaction function (see equation [14]).

As a consequence, both agent B's willingness to perform well and the firm's profits increase.

To examine the total effect, we consider the dependence of firm profits on the strength of altruism n_A as well as the strength of rivalry k_B (Figure 8).⁴² For the sake of simplicity, it is again assumed that the weighting factors in agent A's utility function, n_A and m_A , sum up to 1.0. Figure 8 shows that firm profits are maximized when agent A behaves as altruistically as possible while agent B simultaneously exhibits a well-developed sense of rivalry.

Insert Figure 8 here

Figure 8: The dependence of firm profits on agent A's strength of altruism n_A and agent B's strength of rivalry k_B

This is because agent B's rivalry encourages the reduction of agent A's variable wage payments while at the same time enhancing agent B's own compensation. The altruistic agent A partly internalizes this external effect by drawing

satisfaction from the fact that agent B receives higher remuneration. Thus, the principal's measure to privilege B in order to make use of his or her rivalry simultaneously enhances the utility of altruistic agent A, who attributes to B the improvement of his or her own situation. At the same time, as a result of agent A's altruism, the principal reallocates variable wage compensation components between the decentralized divisions in favor of B and at the expense of A. This action is in agent B's interest whose rivalry prompts him or her to distinguish him- or herself as much as possible from agent A in terms of wage compensation. Therefore, the principal's countermeasures, which aim at balancing the two agents' social preferences, are complementary. The relation between the weighting factors is described by

$$\frac{m_A(1+k_B)+n_Ak_Bl_B}{1+k_B+k_Bl_B} < \frac{m_A(1+k_B)+n_Ak_Bl_B}{m_A-n_A}, \quad (36)$$

meaning that firm profits are in relative terms more enhanced by the agent who behaves competitively. It follows that from a firm perspective, it is advantageous if the altruistic agent is the one who, in the case of completely selfish behavior, would make a smaller contribution to firm profits. Proposition 4 summarizes the previous discussion:

Proposition 4 (Firm Profitability in the Case of Two-Sided Heterogeneous Social Preferences): *For the principal, the optimal combination of types in his or her team is that where the agents are as diverse as possible in terms of social preferences. If additional asymmetries between the agents exist, it is most desirable that the agent who behaves altruistically exhibits higher risk aversion, higher marginal effort costs and a more volatile performance measure.*

This proposition, like the previous ones, relies on the assumptions that separate performance measures for each of the two agents are available and that the principal can observe their (social) preferences.

6. Implications and Further Research

The question of whether new insights gained in experimental decision theory and neuroscience can have an impact on economic theories is an important issue of modern research. In this paper, we analyzed the influence of various social preferences on the incentive system in a decentralized organization, taking into account that in reality not all people behave completely selfishly. Unlike other research in the field, in this paper we also took altruism into account. Thus, this study considers three widespread types of preferences (rivalry, pure self-interest and altruism) and the differences between them. Furthermore, it provides evidence for the importance of internal and external conditions in providing incentives and shaping agent behavior and ultimately firm profits. In particular, we have shown how environmental stochastic dependences, which may be caused by market cycles or other external conditions, influence the principal's provision of incentives. These are crucial to the outcome of our analysis of the impact of social preferences on the profitability of the firm. Our study makes clear that the impact of social preferences interacts with stochastic dependences. As a result, the principal's share parameters fulfill overlapping functions, which may affect the optimality of different types of remuneration. For the sake of simplicity, the influence of internal dependences, in the shape of technological dependences, has been excluded from this paper. In order to optimize the wage compensation system, the principal must react to the agents' preferences by shifting the shares in performance measures that each of the two agents receives and thus reallocating their variable wage compensation components. Our analysis shows that when such a reallocation takes place the different functions that the parameters of the incentive system have must be balanced.

A central result of our paper is the insight that a firm may increase its profit by adjusting compensation. Both altruism and rivalry entail intrinsic motivation, which

can be exploited by the principal, although the altruism of agents is not inherently advantageous. This seems surprising at first glance; however, what makes possible an increase in profits lies in the differences between the preferences of various agents and in the strength of those preferences. A principal, and therefore a firm, can make use of these differences in cases of rivalry as well as of altruism. In the case of rivalry, this is related to the motivating effects of competition, which can be seen in several areas of social interaction, including economics and sports. In the case of altruism, the firm does not need to pay all agents equally in order to motivate them, as it would in a situation where all agents exhibited pure self-interest. Therefore, incentives can be partly shifted toward the selfish or competitive agent who in consequence invests higher effort. This important result indicates that firms, being hierarchically structured organizations, have a chance to exploit intrinsic motivation by combining agents with different social preferences and optimizing the wage compensation system.

The insight that differences in preferences can be profitable and hence can be “managed” by a firm is not only relevant to the issue of wage compensation, but also indicates that ethical aspects should be considered in problems of economics as well as business administration.⁴³ These results are important in determining organizational structure, the distribution of tasks and decision authority within a firm, the selection of the right personnel to fill management positions as well as of the members of business teams including the company’s board. In sum, we could say that since preferences determine human behavior to a high degree, different types of preferences may affect many aspects of a firm.

A significant, and perhaps unexpected, result of our analysis is that a group of people who share equally competitive or altruistic preferences is usually not more efficient than a team composed of members with only selfish preferences and less

efficient than a team composed of members with different preferences. In our model, a firm maximizes its profit when the agents it employs exhibit the greatest possible difference in their (social) preferences. This result raises two questions. First, should firms combine persons with extremely different social preferences in their (project) teams and employ them as managers of decentralized divisions? Second, to what degree can such a policy be realized? Answers to these questions should take the prerequisites of our theoretical model into account, namely (a) that as many performance measures as agents are available and (b) that the principal can observe to some extent his or her agents' (social) preferences. Another point is that our main result conflicts with the intuitive notion, which is supported by some empirical research,⁴⁴ that teams with identical preferences seem to be most successful in specific areas or situations.

All of the aspects discussed above indicate that further research is needed to analyze the relevance and practicability of the assumptions examined here. For example, it seems important that we use a theoretical model without imposing limited liability constraints which is used in conjunction with the assumption that a firm can use very different compensation parameters for each of the two types of agents. The last point contradicts principles of equal treatment in compensation, which calls for empirical research on e.g. the distribution of (social) preferences in companies, their determinants and their stability. The degree to which moral ideas and principles of equality, justice, etc., influence behavior in firms and therefore limit the design of wage compensation systems should also be examined. Additionally, our results need to be scrutinized through empirical investigation to identify in which areas and to what degree firms can use differences in personal preferences to increase their profits.

To conclude we could say that, in reality, people may be driven by a mixture of competitive, selfish and altruistic motives. Assuming that selfish preferences always

prevail is certainly a simplification of real-life behavior. Nevertheless, if we define self-interest as the mean between envy and altruism, which represent extreme types of behavior, this assumption may on average yield satisfactory results. At the same time, this assumption may not be correct in all cases. Behavioral science research has been increasingly revealing differences in people's choices and providing insights into the circumstances under which different preferences are reflected in personal behavior. Such insights should also increasingly be taken into consideration in future theoretical and empirical research in the fields of economics and business administration.

Appendix

A. Derivation of the optimal values for the share rates in a situation which includes rivalry and moral hazard

Solving the participation constraints for α_0 , β_0 and plugging the resulting values together with the expressions for a , b (equation [14]) in the principal's utility function (4) yields his optimization problem:

$$\begin{aligned}
 \max_{\alpha_A, \alpha_B, \beta_A, \beta_B} U_P = & (1+k_A) \frac{1}{c_A} \alpha_A - k_B l_B \frac{1}{c_B} \alpha_B - k_A l_A \frac{1}{c_A} \beta_A + (1+k_B) \frac{1}{c_B} \beta_B - \\
 & \frac{1}{2} \cdot \frac{(1+k_A)^2 (1+k_B + k_B l_B)}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot \frac{1}{c_A} \alpha_A^2 - \frac{1}{2} \cdot \frac{k_B^2 l_B^2 (1+k_A + k_A l_A)}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot \frac{1}{c_B} \alpha_B^2 - \\
 & \frac{1}{2} \cdot \frac{k_A^2 l_A^2 (1+k_B + k_B l_B)}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot \frac{1}{c_A} \beta_A^2 - \frac{1}{2} \cdot \frac{(1+k_B)^2 (1+k_A + k_A l_A)}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot \frac{1}{c_B} \beta_B^2 + \\
 & \frac{k_A l_A (1+k_A)(1+k_B + k_B l_B)}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot \frac{1}{c_A} \alpha_A \beta_A + \frac{k_B l_B (1+k_B)(1+k_A + k_A l_A)}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot \frac{1}{c_B} \alpha_B \beta_B - \\
 & \frac{1+k_B + k_B l_B}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot Risk_A - \frac{1+k_A + k_A l_A}{(1+k_A)(1+k_B) - k_A l_A k_B l_B} \cdot Risk_B,
 \end{aligned} \tag{A.1}$$

where

$$\begin{aligned}
 Risk_A = & \frac{r_A}{2} \cdot \{ [\alpha_A (1+k_A) - k_A l_A \beta_A]^2 \sigma_A^2 + [\alpha_B (1+k_A) - k_A l_A \beta_B]^2 \sigma_B^2 + \\
 & 2 \cdot [\alpha_A (1+k_A) - k_A l_A \beta_A] \cdot [\alpha_B (1+k_A) - k_A l_A \beta_B] \sigma_A \sigma_B \rho \}
 \end{aligned} \tag{A.2}$$

$$\begin{aligned}
 Risk_B = & \frac{r_B}{2} \cdot \{ [\beta_A (1+k_B) - k_B l_B \alpha_A]^2 \sigma_A^2 + [\beta_B (1+k_B) - k_B l_B \alpha_B]^2 \sigma_B^2 + \\
 & 2 \cdot [\beta_A (1+k_B) - k_B l_B \alpha_A] \cdot [\beta_B (1+k_B) - k_B l_B \alpha_B] \sigma_A \sigma_B \rho \}.
 \end{aligned} \tag{A.3}$$

Partial differentiation with respect to the four share rates α_A , β_A , α_B , β_B provides the first-order conditions, which, after some rearranging, can be written as:

$$\text{Condition 1: } \frac{\partial U_P}{\partial \alpha_A} = 0$$

$$\begin{aligned}
 &\Leftrightarrow (1+k_A) [(1+k_A)(1+k_B) - k_A l_A k_B l_B] - \\
 &\left[(1+k_A)^2 (1+k_B + k_B l_B) (1+r_A c_A \sigma_A^2) + k_B^2 l_B^2 (1+k_A + k_A l_A) r_B c_A \sigma_A^2 \right] \alpha_A - \\
 &\left[(1+k_A)^2 (1+k_B + k_B l_B) r_A c_A \sigma_A \sigma_B \rho + k_B^2 l_B^2 (1+k_A + k_A l_A) r_B c_A \sigma_A \sigma_B \rho \right] \alpha_B + \\
 &\left[k_A l_A (1+k_A) (1+k_B + k_B l_B) (1+r_A c_A \sigma_A^2) + \right. \\
 &\quad \left. k_B l_B (1+k_B) (1+k_A + k_A l_A) r_B c_A \sigma_A^2 \right] \beta_A + \\
 &\left[k_A l_A (1+k_A) (1+k_B + k_B l_B) r_A c_A \sigma_A \sigma_B \rho + \right. \\
 &\quad \left. k_B l_B (1+k_B) (1+k_A + k_A l_A) r_B c_A \sigma_A \sigma_B \rho \right] \beta_B = 0
 \end{aligned} \tag{A.4}$$

$$\text{Condition 2: } \frac{\partial U_P}{\partial \alpha_B} = 0$$

$$\begin{aligned}
 &\Leftrightarrow -k_B l_B [(1+k_A)(1+k_B) - k_A l_A k_B l_B] - \\
 &\left[(1+k_A)^2 (1+k_B + k_B l_B) r_A c_B \sigma_A \sigma_B \rho + k_B^2 l_B^2 (1+k_A + k_A l_A) r_B c_B \sigma_A \sigma_B \rho \right] \alpha_A - \\
 &\left[k_B^2 l_B^2 (1+k_A + k_A l_A) (1+r_B c_B \sigma_B^2) + (1+k_A)^2 (1+k_B + k_B l_B) r_A c_B \sigma_B^2 \right] \alpha_B + \\
 &\left[k_A l_A (1+k_A) (1+k_B + k_B l_B) r_A c_B \sigma_A \sigma_B \rho + \right. \\
 &\quad \left. k_B l_B (1+k_B) (1+k_A + k_A l_A) r_B c_B \sigma_A \sigma_B \rho \right] \beta_A + \\
 &\left[k_B l_B (1+k_B) (1+k_A + k_A l_A) (1+r_B c_B \sigma_B^2) + \right. \\
 &\quad \left. k_A l_A (1+k_A) (1+k_B + k_B l_B) r_A c_B \sigma_B^2 \right] \beta_B = 0
 \end{aligned} \tag{A.5}$$

$$\text{Condition 3: } \frac{\partial U_P}{\partial \beta_A} = 0$$

$$\begin{aligned}
 &\Leftrightarrow -k_A l_A [(1+k_A)(1+k_B) - k_A l_A k_B l_B] + \\
 &\left[k_A l_A (1+k_A) (1+k_B + k_B l_B) (1+r_A c_A \sigma_A^2) + \right. \\
 &\quad \left. k_B l_B (1+k_B) (1+k_A + k_A l_A) r_B c_A \sigma_A^2 \right] \alpha_A + \\
 &\left[k_A l_A (1+k_A) (1+k_B + k_B l_B) r_A c_A \sigma_A \sigma_B \rho + \right. \\
 &\quad \left. k_B l_B (1+k_B) (1+k_A + k_A l_A) r_B c_A \sigma_A \sigma_B \rho \right] \alpha_B - \\
 &\left[k_A^2 l_A^2 (1+k_B + k_B l_B) (1+r_A c_A \sigma_A^2) + (1+k_B)^2 (1+k_A + k_A l_A) r_B c_A \sigma_B^2 \right] \beta_A - \\
 &\left[k_A^2 l_A^2 (1+k_B + k_B l_B) r_A c_A \sigma_A \sigma_B \rho + (1+k_B)^2 (1+k_A + k_A l_A) r_B c_A \sigma_A \sigma_B \rho \right] \beta_B = 0
 \end{aligned} \tag{A.6}$$

$$\text{Condition 4: } \frac{\partial U_P}{\partial \beta_B} = 0$$

$$\begin{aligned}
 &\Leftrightarrow (1+k_B) [(1+k_A)(1+k_B) - k_A l_A k_B l_B] + \\
 &\left[k_A l_A (1+k_A) (1+k_B + k_B l_B) r_A c_B \sigma_A \sigma_B \rho + \right. \\
 &\quad \left. k_B l_B (1+k_B) (1+k_A + k_A l_A) r_B c_B \sigma_A \sigma_B \rho \right] \alpha_A + \\
 &\left[k_B l_B (1+k_B) (1+k_A + k_A l_A) (1+r_B c_B \sigma_B^2) + \right. \\
 &\quad \left. k_A l_A (1+k_A) (1+k_B + k_B l_B) r_A c_B \sigma_B^2 \right] \alpha_B - \\
 &\left[k_A^2 l_A^2 (1+k_B + k_B l_B) r_A c_B \sigma_A \sigma_B \rho + (1+k_B)^2 (1+k_A + k_A l_A) r_B c_B \sigma_A \sigma_B \rho \right] \beta_A - \\
 &\left[(1+k_B)^2 (1+k_A + k_A l_A) (r_B^2 + r_B c_B \sigma_B^2) + k_A^2 l_A^2 (1+k_B + k_B l_B) r_A c_B \sigma_B^2 \right] \beta_B = 0
 \end{aligned} \tag{A.7}$$

Having rearranged and simplified the first-order conditions to some extent, solved

equation (A.4) for α_B and plugged the resulting expression into equation (A.6) we get:

$$\alpha_A = \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{(1+k_A)(1+k_B + k_B l_B)} - \frac{1+k_B}{k_B l_B} \cdot r_A c_A \sigma_A^2 \cdot \beta_A + \frac{k_A l_A}{1+k_A} \cdot (1+r_A c_A \sigma_A^2) \cdot \beta_A - \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{(1+k_A)k_B l_B} \cdot r_A c_A \sigma_A \sigma_B \rho \cdot \beta_B. \quad (A.8)$$

Solving equation (A.5) for α_B and plugging the resulting expression into equation (A.7) yields:

$$\alpha_A = \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{k_B l_B (1+k_A + k_A l_A)} \cdot \frac{\sigma_B}{\sigma_A \rho} - \frac{1+k_B}{k_B l_B} \cdot r_B c_B \sigma_B^2 \cdot \beta_A + \frac{k_A l_A}{1+k_A} \cdot (1+r_B c_B \sigma_B^2) \cdot \beta_A - \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{k_B l_B (1+k_A)} \cdot (1+r_B c_B \sigma_B^2) \cdot \frac{\sigma_B}{\sigma_A \rho} \beta_B. \quad (A.9)$$

After rearranging and simplification, solving equation (A.4) for α_A and plugging the resulting expression into equation (A.6) leads to:

$$\alpha_B = - \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{(1+k_A)(1+k_B + k_B l_B)} \cdot \frac{\sigma_A}{\sigma_B \rho} + \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{k_B l_B (1+k_A)} \cdot (1+r_A c_A \sigma_A^2) \cdot \frac{\sigma_A}{\sigma_B \rho} \beta_A + \frac{1+k_B}{k_B l_B} \cdot (1+r_A c_A \sigma_A^2) \cdot \beta_B - \frac{k_A l_A}{1+k_A} \cdot r_A c_A \sigma_A^2 \cdot \beta_B. \quad (A.10)$$

Solving equation (A.5) for α_A and plugging the resulting expression into equation (A.7) yields in its turn:

$$\alpha_B = - \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{k_B l_B (1+k_A + k_A l_A)} + \frac{[(1+k_A)(1+k_B) - k_A l_A k_B l_B]}{k_B l_B (1+k_A)} \cdot \frac{r_B c_B \sigma_A \sigma_B \rho}{p_B^2} \beta_A + \frac{1+k_B}{k_B l_B} \cdot \frac{(p_B^2 + r_B c_B \sigma_B^2)}{p_B^2} \beta_B - \frac{k_A l_A}{1+k_A} \cdot \frac{r_B c_B \sigma_B^2}{p_B^2} \beta_B. \quad (A.11)$$

Equating the expressions for α_A in (A.8) and (A.9) leads to the mathematical relation:

$$\frac{1+k_A}{1+k_A + k_A l_A} \cdot \frac{\sigma_B}{\sigma_A \rho} - \frac{k_B l_B}{1+k_B + k_B l_B} - (r_B c_B \sigma_B^2 - r_A c_A \sigma_A^2) \beta_A - \left[(1+r_B c_B \sigma_B^2) \cdot \frac{\sigma_B}{\sigma_A \rho} - r_A c_A \sigma_A \sigma_B \rho \right] \beta_B = 0. \quad (A.12)$$

Accordingly, equating the expressions for α_B in (A.10) and (A.11) yields:

$$\begin{aligned}
 & - \frac{1+k_A}{1+k_A+k_A l_A} + \frac{k_B l_B}{1+k_B+k_B l_B} \cdot \frac{\sigma_A}{\sigma_B \rho} + (r_B c_B \sigma_B^2 - r_A c_A \sigma_A^2) \beta_B + \\
 & \left[r_B c_B \sigma_A \sigma_B \rho - (1+r_A c_A \sigma_A^2) \cdot \frac{\sigma_A}{\sigma_B \rho} \right] \beta_A = 0.
 \end{aligned} \tag{A.13}$$

Solving (A.12) for β_A and plugging the resulting expression into (A.13) delivers, after some further rearranging and simplification, the optimal value of share rate β_B :

$$\begin{aligned}
 \beta_B = & \frac{1+k_A}{1+k_A+k_A l_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} - \\
 & \frac{k_B l_B}{1+k_B+k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \cdot \frac{\sigma_A}{\sigma_B} \rho.
 \end{aligned} \tag{A.14}$$

Plugging this expression into (A.12) leads to the optimal value of share rate β_A :

$$\begin{aligned}
 \beta_A = & \frac{k_B l_B}{1+k_B+k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} - \\
 & \frac{1+k_A}{1+k_A+k_A l_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} \cdot \frac{\sigma_B}{\sigma_A} \rho.
 \end{aligned} \tag{A.15}$$

Substituting the expressions (A.14) and (A.15) for β_B and β_A in the mathematical relations (A.8) and (A.11) delivers the optimal values for the share rates α_A and α_B :

$$\begin{aligned}
 \alpha_A = & \frac{1+k_B}{1+k_B+k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} - \\
 & \frac{k_A l_A}{1+k_A+k_A l_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} \cdot \frac{\sigma_B}{\sigma_A} \rho
 \end{aligned} \tag{A.16}$$

$$\begin{aligned}
 \alpha_B = & \frac{k_A l_A}{1+k_A+k_A l_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} - \\
 & \frac{1+k_B}{1+k_B+k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \cdot \frac{\sigma_A}{\sigma_B} \rho.
 \end{aligned} \tag{A.17}$$

B. Calculation of the principal's optimal goal function value in a situation of rivalry

We calculate the maximum value of the objective function by plugging the optimal values for the share rates (see equations [17]–[22]) into the principal's goal function (A.1). The expressions in front of the risk terms yield:

$$\begin{aligned}
 & \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \frac{1}{c_A} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} - \\
 & \frac{1}{2} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \frac{1}{c_A} \cdot \left[\frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \right]^2 + \\
 & \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_A+k_A l_A} \cdot \frac{1}{c_B} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} - \\
 & \frac{1}{2} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_A+k_A l_A} \cdot \frac{1}{c_B} \cdot \left[\frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} \right]^2.
 \end{aligned} \tag{A.18}$$

Agent A's risk term can be simplified to:

$$\begin{aligned}
 & - \frac{r_A \sigma_A^2}{2} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \left[\frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \right]^2 - \\
 & \frac{r_A}{2} \cdot \sigma_A^2 \rho^2 \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \left[\frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \right]^2 + \\
 & r_A \sigma_A^2 \rho^2 \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \left[\frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \right]^2 = \\
 & - \frac{r_A \sigma_A^2 (1-\rho^2)}{2} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \left[\frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} \right]^2.
 \end{aligned} \tag{A.19}$$

Analogously, one obtains for agent B's risk term:

$$- \frac{r_B \sigma_B^2 (1-\rho^2)}{2} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_A+k_A l_A} \cdot \left[\frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} \right]^2. \tag{A.20}$$

Consolidating the partial expressions (A.18)–(A.20), after some rearranging yields:

$$\begin{aligned}
 U_P = & \frac{1}{c_A} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} - \\
 & \frac{1}{2c_A} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_B+k_B l_B} \cdot \frac{1+r_A \sigma_A^2 (1-\rho^2) c_A}{\left[1+r_A \sigma_A^2 (1-\rho^2) c_A \right]^2} + \\
 & \frac{1}{c_B} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_A+k_A l_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B} - \\
 & \frac{1}{2c_B} \cdot \frac{\left[(1+k_A)(1+k_B) - k_A l_A k_B l_B \right]}{1+k_A+k_A l_A} \cdot \frac{1+r_B \sigma_B^2 (1-\rho^2) c_B}{\left[1+r_B \sigma_B^2 (1-\rho^2) c_B \right]^2}.
 \end{aligned} \tag{A.21}$$

After some further manipulations, one finally obtains the expression given in the main text:

$$\begin{aligned}
 U_P^*(k_A, k_B) = & \frac{1}{2c_A} \cdot \frac{(1+k_A)(1+k_B) - k_A l_A k_B l_B}{1+k_B+k_B l_B} \cdot \frac{1}{1+r_A \sigma_A^2 (1-\rho^2) c_A} + \\
 & \frac{1}{2c_B} \cdot \frac{(1+k_A)(1+k_B) - k_A l_A k_B l_B}{1+k_A+k_A l_A} \cdot \frac{1}{1+r_B \sigma_B^2 (1-\rho^2) c_B}.
 \end{aligned}
 \tag{A.22}$$

References

Andreoni, J., Harbaugh, W.T. and Vesterlund, L., 2007, "Altruism in Experiments," *Working Paper, Universities of California, Oregon and Pittsburgh*.

Andreoni, J. and Miller, J., 2002, "Giving According to GARP. An Experimental Test of the Consistency of Preferences for Altruism," *Econometrica*, 70, 737-753.

Banker, R.D. and Datar, S.M., 1989, "Sensitivity, Precision, and Linear Aggregation of Signals for Performance Evaluation," *Journal of Accounting Research*, 27, 21-39.

Bartling, B. and von Siemens, F., 2004, "Efficiency in Team Production with Inequity Averse Agents," *Working Paper, University of Munich*.

Bartling, B. and von Siemens, F., 2005, "Inequity Aversion and Moral Hazard with Multiple Agents," *Working Paper, University of Munich*.

Bartling, B. and von Siemens, F., 2006, "The Intensity of Incentives in Firms and Markets: Moral Hazard with Envious Agents," *Working Paper, University of Munich*.

Bester, H. and Güth, W., 1998, "Is Altruism Evolutionarily Stable?," *Journal of Economic Behavior and Organization*, 34, 193–209.

Camerer, C.F., 2003, *Behavioral Game Theory*, Princeton: Princeton University Press.

Camerer, C.F., Loewenstein, G. and Prelec, D., 2004, "Neuroeconomics. Why Economics Needs Brains," *Scandinavian Journal of Economics*, 106, 555-579.

Christensen, P.O. and Feltham, G.A., 2005, *Economics of Accounting. Performance Evaluation*, New York: Springer.

Datar, S.M., Kulp, S.C. and Lambert, R.A., 2001, "Balancing Performance Measures," *Journal of Accounting Research*, 39, 75-92.

Demougin, D. and Fluet, C., 2006, "Group vs. Individual Performance Pay when Workers are Envious," D. Demougin and C. Schade, eds. *Contributions to Entrepreneurship and Economics. First Haniel-Kreis Meeting on Entrepreneurial Research*, Berlin: Duncker and Humblot, (forthcoming).

Demougin, D., Fluet, C. and Helm, C., 2006, "Output and Wages with Inequity Averse Agents," *Canadian Journal of Economics*, 39, 399-413.

Dierkes, S. and Harreiter, B., 2006, "Soziale Präferenzen und relative Leistungsbewertung. Eine agency-theoretische Analyse," *Working Paper*.

Dur, R. and Glazer, A., 2008, "Optimal Incentive Contracts when a Worker Envy his Boss," *Journal of Law, Economics and Organization*, (forthcoming).

Edgeworth, F.Y., 1881, *Mathematical Physics. An Essay on the Application of Mathematics to the Moral Sciences*, London: Kegan Paul & Co.

Englmaier, F. and Wambach, A., 2005, "Optimal Incentive Contracts under Inequity Aversion," *IZA Working Paper No. 1643, University of Munich and IZA Bonn*.

Ewert, R. and Wagenhofer, A., 1993, "Linearität und Optimalität in ökonomischen Agency Modellen. Zur Rechtfertigung des LEN-Modells," *Zeitschrift für Betriebswirtschaft*, 63, 373-391.

Fehr, E. and Falk, A., 2002, "Psychological Foundations of Incentives," *European Economic Review*, 46, 687-724.

Fehr, E. and Fischbacher, U., 2002, "Why Social Preferences Matter. The Impact of Non-Selfish Motives on Competition, Cooperation and Incentives," *The Economic Journal*, 112, C1-C33.

Fehr, E., Fischbacher, U. and Kosfeld, M., 2005, "Neuroeconomic Foundations of Trust and Social Preferences. Initial Evidence," *American Economic Review*, 95, 346-351.

Fehr, E. and Schmidt, K.M., 1999, "A Theory of Fairness, Competition and Cooperation," *Quarterly Journal of Economics*, 114, 817-868.

Fehr, E. and Schmidt, K.M., 2003, "Theories of Fairness and Reciprocity. Evidence and Economic Applications," M. Dewatripont, L.P. Hansen and S.J. Turnovsky, eds. *Advances in Economics and Econometrics*, 8th World Congress, Cambridge: Econometric Society Monographs, 208-257.

Feltham, G.A. and Xie, J., 1994, "Performance Measure Congruity and Diversity in Multi-Task Principal / Agent Relations," *The Accounting Review*, 69, 429-453.

Financial Times Deutschland, 2007, "*Management Salaries. A Question of Performance*," Unknown Author, Hamburg, December 14th 2007, p. 31.

Glimcher, P.W., 2003, *Decisions, Uncertainty and the Brain. The Science of Neuroeconomics*, Cambridge, London: The MIT Press.

Grossmann, K., 2007, "The duty of board of directors," *Financial Times Deutschland*, December 12th 2007, 33.

Grund, C. and Sliwka, D., 2005, "Envy and Compassion in Tournaments," *Journal of Economics and Management Strategy*, 14, 187-207.

Hemmer, T., 2004, "Lessons Lost in Linearity: A Critical Assessment of the General Usefulness of LEN Models in Compensation Research," *Journal of Management Accounting Research*, 16, 149-162.

Henrich, J., 2004, "Cultural Group Selection, Coevolutionary Processes and Large-Scale Cooperation," *Journal of Economic Behavior and Organization*, 53, 3-35.

Hobman, E.V., Bordia, P. and Gallois C., 2003, "Consequences of Feeling Dissimilar From Others in a Work Team," *Journal of Business and Psychology*, 17, 301-325.

Holmström, B., 1979, "Moral Hazard and Observability," *Bell Journal of Economics*, RAND 10, 74-91.

Holmström, B. 1982, "Moral Hazard in Teams," *Bell Journal of Economics*, RAND 13, 324-340.

Holmström, B. and Milgrom, P.R., 1987, "Aggregation and Linearity in the Provision of Intertemporal Incentives," *Econometrica*, 55, 303-328.

Huck, S., Kübler, D. and Weibull, J., 2006, "Social Norms and Optimal Incentives in Firms," *Working Paper, ERSC Center for Economic Learning and Social Evolution*.

Hulverscheidt, C., 2008, "Non-Justifiable Remuneration," *Süddeutsche Zeitung*, 158, Wednesday, July 9th 2008, 19.

Itoh, H., 2004, "Moral Hazard and Other-Regarding Preferences," *Japanese Economic Review*, 55, 18-45.

Küpper, H.-U., 2006, *Unternehmensethik. Hintergründe, Konzepte, Anwendungsbereiche*, Stuttgart: Schäffer-Poeschel.

Lévy-Garboua, L., Meidinger, C. and Rapoport, B., 2008, "The Formation of Social Preferences. Some Lessons from Psychology and Biology," L. A. Gérard-Varet, S.-C. Kolm and J. Mercier-Ythier, eds. *A Handbook on the Economics of Giving, Reciprocity and Altruism*, Amsterdam: North-Holland, (forthcoming).

Mayer, B., 2006, "Prinzipien der Anreizgestaltung bei sozialen Interaktionen. Eine theoriegeleitete Analyse," *Unpublished Dissertation*.

Mayer, B. and Pfeiffer, Th., 2004, "Prinzipien der Anreizgestaltung bei Risikoaversion und sozialen Präferenzen," *Zeitschrift für Betriebswirtschaft*, 74, 1047-1075.

Milgrom P.R. and Roberts, J., 1992, *Economics, Organization and Management*, New Jersey: Prentice Hall.

Neilson W.S. and Stowe, J., 2008, "Piece Rate Contracts for Other-Regarding Workers," *Economic Inquiry*, (forthcoming).

Rey Biel, P., 2007, "Inequity Aversion and Team Incentives," *Working Paper, University College London*.

Rob, R. and Zemsky, P., 2002, "Cooperation, Corporate Culture and Incentive Intensity," *Bell Journal of Economics*, RAND 33, 243-257.

Sandner K.J., 2008a, *Behavioral Contract Theory. Einfluss sozialer Präferenzen auf die Steuerung dezentraler Organisationseinheiten*, Berlin: Duncker and Humblot.

Sandner K.J., 2008b, "Impacts of Rivalry on Types of Compensation. Competition vs. Co-operation amongst Multiple Agents under Technological Interdependencies," *Zeitschrift für Betriebswirtschaft*, forthcoming.

Spremann, K., 1987, "Agent and Principal," G. Bamberg and K. Spremann, eds. *Agency Theory, Information, and Incentives*, Berlin: Springer, 3-37.

The Economist, 2006, "*The Rich, the Poor and the Growing Gap Between Them*," Unknown Author, London, June 17th 2006, 28-30.

The Economist, 2007, "*Rich Man, Poor Man*," Unknown Author, London, January 20th 2007, 11-12.

¹ See Holmström (1979, 1982) for basic principles and Milgrom/Roberts (1992) for an overview.

² See Henrich (2004); Glimcher (2003); Lévy-Garboua, Meidinger and Rapoport (2008).

³ See Camerer, Loewenstein, and Prelec (2004), p. 570ff.; also Fehr, Fischbacher and Kosfeld (2005).

⁴ See Fehr and Falk (2002); Fehr and Fischbacher (2002); also Fehr and Schmidt (2003) as well as Camerer (2003), pp. 43–113.

⁵ See Fehr and Fischbacher (2002), pp. C2-C4 as well as Fehr and Schmidt (2003), p. 218ff.

⁶ Fehr and Fischbacher (2002), C2.

⁷ See Fehr and Fischbacher (2002), pp. C2-C4.

⁸ See Grossmann (2007); the article "Management Salaries. A Question of Performance" in the newspaper "Financial Times Deutschland" (2007) and Hulverscheidt (2008).

⁹ See e.g. the articles “Poor Man, Rich Man” as well as “The Rich, the Poor and the Growing Gap Between Them” in the newspaper “The Economist” (2006, 2007).

¹⁰ See Andreoni (2007) for an overview of experimental evidence so far.

¹¹ Fairness-based preferences, in contrast, would have to be modeled in quite a different way (see Fehr and Schmidt 1999, 2003). We therefore excluded these types of social preferences from our study.

¹² See Itoh (2004); Mayer and Pfeiffer (2004); Englmaier and Wambach (2005); Mayer (2006); Dur and Glazer (2008).

¹³ See Itoh (2004); Bartling and von Siemens (2004, 2005, 2006); Demougin and Fluet (2006); Demougin, Fluet and Helm (2006); Rey Biel (2007); Neilson and Stowe (2008).

¹⁴ See, for example, Itoh (2004); Bartling and von Siemens (2004, 2005, 2006); Demougin, Fluet and Helm (2006); Rey Biel (2007).

¹⁵ The findings of Bartling and von Siemens (2004); Grund and Sliwka (2005); Rob and Zemsky (2002); but also Huck, Kübler and Weibull (2006) bear some relevance to our study.

¹⁶ See, for example, the model specifications used by Englmaier and Wambach (2005); Bartling and von Siemens (2005), as well as Demougin and Fluet (2006).

¹⁷ Itoh (2004) examines a discrete two-point model that can be solved for the principal's decision parameters which specify the dependence of the agents' remuneration on observed outcomes.

¹⁸ See Itoh (2004).

¹⁹ For early applications of the LEN model see Spremann (1987), as well as Holmström and Milgrom (1987). Assessments of its applicability based on the underlying linearity assumptions are given in Ewert and Wagenhofer (1993) as well as in Hemmer (2004).

²⁰ For research in the case of purely selfish preferences, see Banker and Datar (1989); Feltham and Xie (1994), as well as Datar, Kulp and Lambert (2001). For an overview, see Christensen and Feltham (2005). Mayer and Pfeiffer (2004) analyze the impact of rivalry on the balancing of performance measures in a context where the principal has to motivate only one agent, whose wage payments the principal can make contingent on firm profits and one or more additional performance measures.

²¹ See, for example, Itoh (2004), as well as Bartling and von Siemens (2006).

²² Most of the studies concerned with horizontal social preferences assume, for example, risk-neutrality. See Itoh (2004); Bartling and von Siemens (2004, 2005); Demougin and Fluet (2006); Rey Biel (2007); Neilson and Stowe (2008). By contrast, we allow for varying degrees of risk aversion. Furthermore, differences in the agents' (marginal) effort costs are explicitly taken into account.

²³ For a formal analysis of the optimal contract design, taking into account social preferences and technological as well as stochastic interdependences, see Sandner (2008b).

²⁴ See Sandner (2008a), p. 81ff.

²⁵ This specification has also been employed by Andreoni and Miller (2002), who used it to explain altruistic behavior for 20% of the participants in a dictator game.

²⁶ The exposition in (9) corresponds to the one in Bester and Güth (1998) and goes back to Edgeworth (1881). In this case, the parameter values n_i can only lie in the intervals $n_i \in [0; 1]$.

²⁷ For the relevant calculations, see the Appendix.

²⁸ The principal's decision problem in the situation of purely selfish behavior under stochastic dependences can be broken down into two problems; one for agent A and one for agent B. In a situation with differing (social) preferences but uncorrelated error terms ($\rho = 0$), it can again be broken down, but this time into one problem for performance measure x_A and another for x_B . The simultaneous consideration of psychological as well as stochastic dependences means that each share rate fulfills several functions at a time. As a consequence, dividing the principal's optimization problem further is not possible.

²⁹ Since β_A can only take negative values, agent A's risk premium always increases when the absolute value of β_A increases.

³⁰ The subscript for k_A is suppressed in the figure.

³¹ Concerning the exogenous parameters in expression (27) for share rate α_B , values $c_A = c_B = r_A = r_B = \sigma_A^2 = \sigma_B^2 = 1, l_A = 4$ were chosen. All major insights gained from this can be generalized to any other configuration of parameters.

³² Note that β_B is strictly positive. It is the only factor that enlarges agent B's wage payments and that thus becomes negatively reflected in agent A's social preference.

³³ See agent A's reaction function in equation (15).

³⁴ See the Appendix for the relevant calculations.

³⁵ The parameter values used for the underlying simulation are $c_A = c_B = r_A = r_B = \sigma_A^2 = \sigma_B^2 = l_A = 1$.

³⁶ For high values of the aspiration level l_A and simultaneously small values of k_A , there exists the possibility that rivalry is actually detrimental. However, this result is a rather special case and therefore of primarily theoretical interest.

³⁷ The underlying parameter configuration is $c_A = r_A = \sigma_A = 4, c_B = r_B = \sigma_B = 0.5, \rho = 0.5, l_A = 1$.

³⁸ Note that all other parameters have to be kept constant for calculating the threshold values.

³⁹ See Itoh (2004), p. 42 for a similar result. Differences are described in the considerations below.

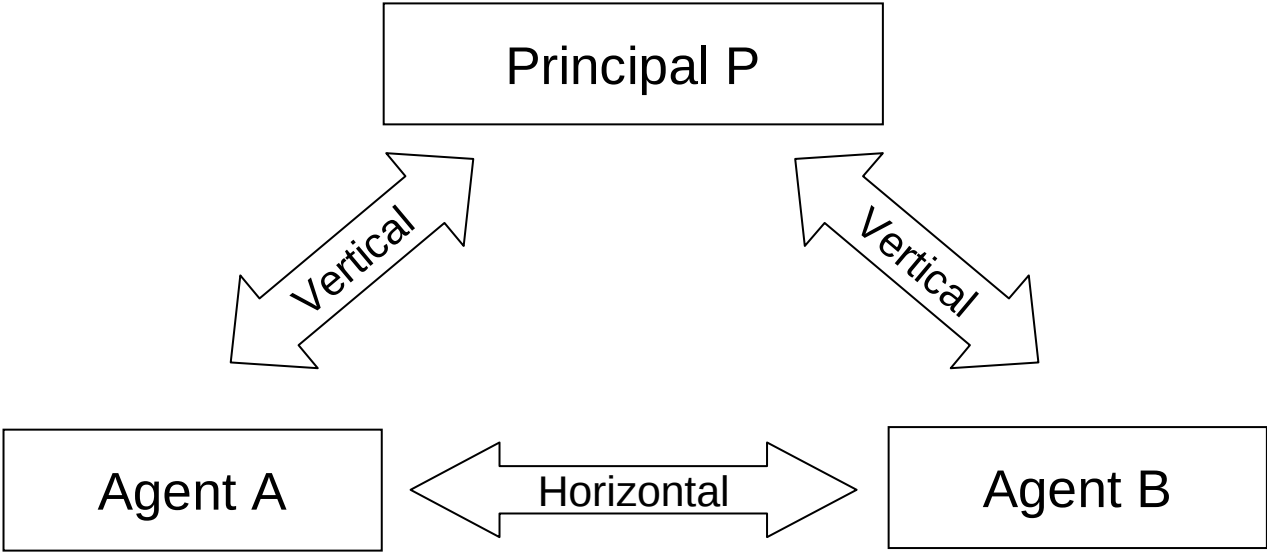
⁴⁰ The parameter configuration is $c_A = c_B = r_A = r_B = \sigma_A^2 = \sigma_B^2 = l_A = l_B = 1, \rho = 0.5$. In the case of altruism, it is additionally assumed that m_A and n_A as well as m_B and n_B sum up to 1.

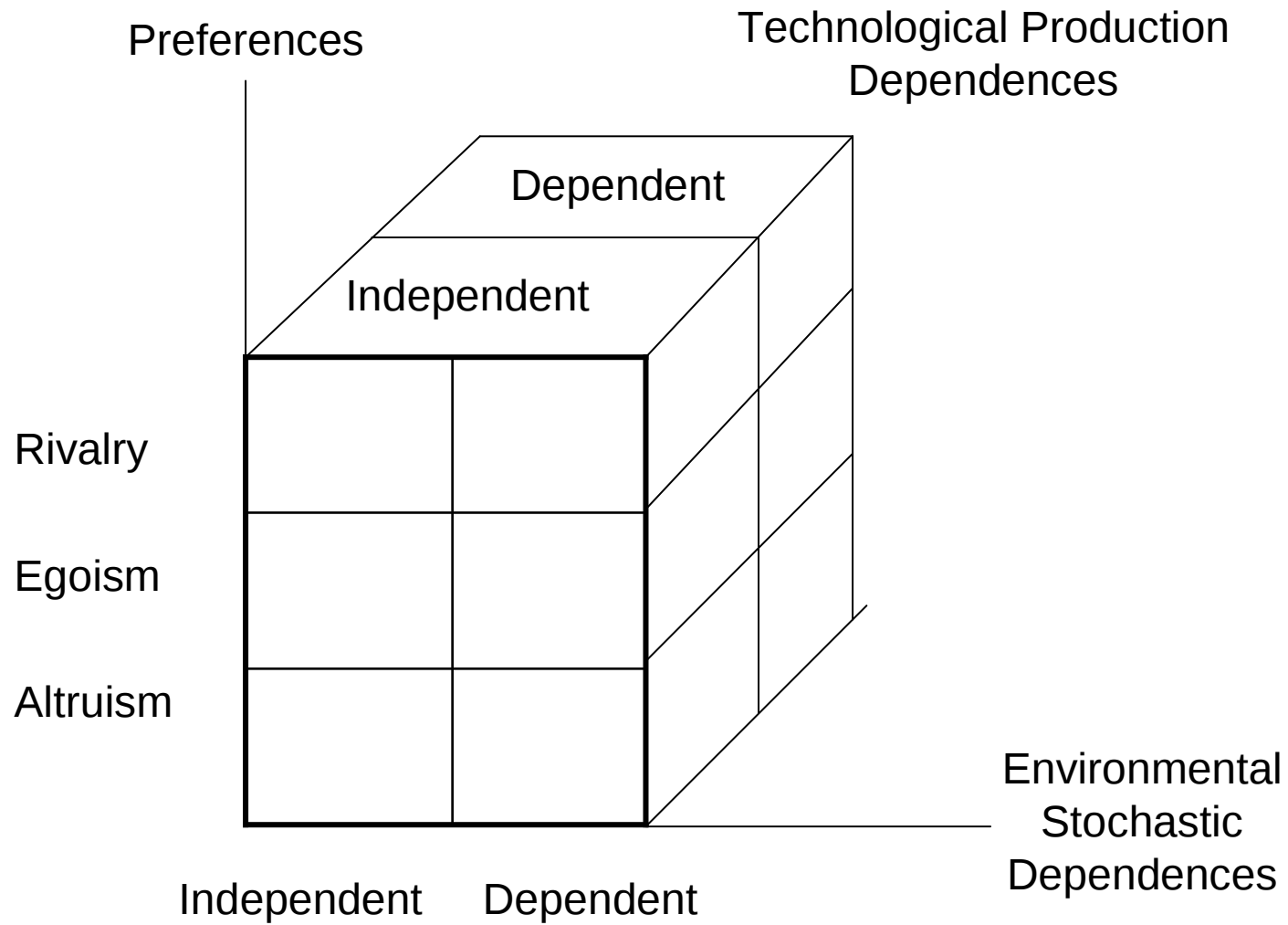
⁴¹ This contradicts the insight of Itoh (2004), who concludes that the principal can make use of his or her agents' competitive behavior. Itoh bases these conclusions on a model with completely identical agents.

⁴² The parameter configuration is $c_A = c_B = r_A = r_B = \sigma_A^2 = \sigma_B^2 = l_B = 1, \rho = 0.5$.

⁴³ For the conceptual foundations of business ethics from a business-economist's point of view see also Küpper (2006).

⁴⁴ See e.g. Hobman, Bordia and Gallois (2003).





Agent A \ Agent B	Rivalry	Egoism	Altruism
Rivalry	R R	R E	R A
Egoism	E R	E E	E A
Altruism	A R	A E	A A

