



universität
wien

MAGISTERARBEIT

Titel der Magisterarbeit

Social Preferences in Delegated Bargaining

Verfasserin

Nicole Fromvald, Bakk.

angestrebter akademischer Grad

Magistra der Sozial- und Wirtschaftswissenschaften

(Mag. rer. soc. oec.)

Wien, im Februar 2013

Studienkennzahl lt. Studienblatt:

A 066 915

Studienrichtung lt. Studienblatt:

Magisterstudium Betriebswirtschaft

Betreuer:

Univ.-Prof. Dr. Thomas Pfeiffer

Danksagung

An dieser Stelle bedanke ich mich bei allen Personen,
die mich im Laufe meines Studiums und während der Verfassung
meiner Magisterarbeit unterstützt haben.

Besonderer Dank gilt vor allem meiner **Mama**, die mir nicht nur
während der gesamten Studienzeit zur Seite stand
sondern mich in allen Lebenssituationen
voll unterstützt.

Für die fachliche Unterstützung bei der Erstellung dieser Arbeit
danke ich meinem Betreuer Univ.-Prof. Dr. Thomas Pfeiffer
sowie Dr. Clemens Löffler für ihren wertvollen fachlichen Rat
und die hilfreichen Anregungen.

List of Contents

List of Tables	IV
List of Figures.....	V
List of Variables and Design Parameters	VI
List of Abbreviations	VII
1 Introduction	1
1.1 Problem statement.....	1
1.2 Constitution	4
2 The nature of social preferences.....	4
2.1 Reciprocity	5
2.2 Fairness – Inequity aversion	5
2.3 Altruism	6
3 Overview of models considering social preferences	7
4 Description of the models.....	11
4.1.1 Introduction.....	11
4.1.2 The models	12
5 Individual work without interaction.....	19
5.1 Fairness in delegated bargaining	20

5.1.1	Choice of effort.....	20
5.1.2	Optimal employment contracts.....	21
5.1.3	Choice of agent.....	23
5.2	Moral Hazard and other-regarding preferences.....	23
5.2.1	Self-interested agent.....	23
5.2.2	Other-regarding agent.....	25
5.2.3	Wage if project fails is 0	26
5.2.4	Wage if project succeeds is higher than 0.....	28
5.2.5	Choice of agent.....	29
5.2.6	“Net” Material Payoff	31
5.3	Comparison of “No Interaction”.....	33
6	Interaction with a third party.....	36
6.1	Fairness in delegated bargaining.....	36
6.1.1	Interaction with an external third party	36
6.1.1.1	Introduction.....	36
6.1.1.2	Scaling factor: Third party’s payoff less sensitive to agent’s effort..	36
6.1.2	Interaction with an internal third party	37
6.1.2.1	The model.....	37

6.1.2.2	Choice of effort.....	39
6.1.2.3	Optimal employment contracts.....	41
6.1.2.4	Choice of agent.....	44
6.1.2.4.1	Same concern for both.....	45
6.1.2.4.2	No concern for principal but concern for third party	45
6.1.2.4.3	Influence of the spillover factor	47
6.2	Moral hazard and other-regarding agents	48
6.2.1	The model.....	48
6.2.2	The self-interest case	50
6.2.3	Other-regarding agents.....	51
6.2.4	“Net” Material Payoff.....	55
6.3	Comparison of “Interaction”.....	58
7	Conclusion	62
8	Summary.....	65
8.1	English	65
8.2	German	66
	Bibliography.....	67
	Curriculum vitae	71

List of Tables

Table 1: The settings 14

Table 2: Comparison of "No Interaction" 33

Table 3: Comparison of "Interaction" 58

List of Figures

Figure 1: Utility function Inequity Aversion8

Figure 2: Principal-Agent-Model with fairness considerations13

Figure 3: Convex utility function15

Figure 4: Timing17

Figure 5: Individual Work without Interaction.....19

Figure 6: Only concern for third party46

Figure 7: Choice of agent in the internal setting47

List of Variables and Design Parameters

π_P	Expected payoff of principal
π_A	Expected payoff of agent
π_B	Expected payoff of third party
U_S	Utility function of selfish agent
U_F	Utility function of fair agent
u_A	Agent's utility function
u_P	Principal's utility function
μ_P	Concern factor for principal
μ_B	Concern factor for third party
γ	Inequity aversion
θ_L	Low output
θ_H	High output
e	Effort
$e(S)$	Effort of selfish agent
$e(F)$	Effort of fair agent
$C(e)$	Cost of effort
$q_A(e)$	Output produced by agent
$q_B(e)$	Output produced by third party
$q'_A(e)$	Marginal output produced by agent
$q'_B(e)$	Marginal output produced by third party
$E(w)$	Agent's expected wage payment
α	Fixed payment
β	Agent's commission rate
r	Price
h	Scaling factor
γ	Third party's commission rate
δ	Fixed payment of third party
p	Probability
w	Wage
$\bar{u}$	Reservation utility
∂	Differentiation
Δ	Delta (change)

List of Abbreviations

e.g. for example

s.t. subject to

Social preferences in delegated bargaining

1 Introduction

1.1 Problem statement

“Models of social preferences assume, in contrast, that the decision maker may also care about how much material resources are allocated to others”.¹

According to the standard economic, neoclassical theory of the firm, the actors on the market are seen as homo oeconomicus, profit-maximizing individuals². The assumption, that people’s choices about actions are only concerned about their own interest was the basis for many models starting for example with Adam Smith.³ This assumption is of course a convenient simplification of the real world situation.⁴

In recent years the economic literature has emerged social preferences of actors. So not all individuals are strictly selfish, there exist also people that care about the wellbeing or the payoff of others.⁵ Such fairness preferences concern not only friends, family, colleagues, but even casual strangers⁶. So the standard economy can be refuted, as the assumption that all people are selfish is not right

¹ Fehr, E., Schmidt, K. (2002), page 12

² See Aoki, M. (1984), page 11f

³ See Frohlich et al. (2004), page 91

⁴ See Fehr, E., Fischbacher, U. (2002), page 29

⁵ See Lammers, F. (2010), page 1f

⁶ See Frohlich et al. (2004), page 91

“anymore”.⁷ These so-called behavioral economics use psychological regularity so as to diminish rationality assumptions and extend the theory⁸. Of course there are people with other-regarding, social, fair preferences as data from experiments on ultimatum games, public good games, gift exchange games, and others demonstrate⁹.

The purpose of this work is to show that economists fail if they persist on the self-interest hypothesis and rule out any social preferences as many experiments from the last decades show that social preferences indeed matter.¹⁰ But in fact, even if many models deal with social preferences, either reciprocity, fairness, inequity aversion, and so forth, there is no clear outcome observable.

Generally, there exist people concerning about other’s wellbeing, so acting fairly. In contrast selfish people are only concerned about their own payoff.¹¹ Both types of agents then have a different impact on the firm’s wellbeing as will be described later. This conclusion, whether selfish or fair agents are preferred, depends in fact on the individual assumptions made. In order to show that solutions vary depending on the individual assumptions, formulas and situations (individual work or interaction with another agent), the following two models are presented and discussed:

- Lammers, F. (2010): “Fairness in Delegated Bargaining”¹²
This article is chosen in order to make a first introduction into this topic as it easily tries to explain this theory with the help of models illustrating the complexity of social preferences. Secondly it shows that different assumptions result in different solutions compared to other models. This

⁷ See Fehr, E., Fischbacher, U. (2002), page 1

⁸ See Camerer, C.F. (2003), page 3

⁹ See Itoh, H. (2004), page 18

¹⁰ See Fehr, E., Fischbacher, U. (2002), page 1

¹¹ See Fehr, E., Schmidt, K.M. (1999), page 817-820

¹²¹² See Lammers, F. (2010), page 169ff

model is analyzed under different situations, on the one hand the agent works individually and on the other hand he collaborates with a third party.

- Itoh, H. (2004): "Moral hazard and other-regarding preferences"¹³
The article is chosen as it equals the above article in many (not all!) assumptions. Especially the situations, the agent is working, is equal but points out different functions. It also tries to explain the complexity of social preferences with mathematical models. In fact it shows that under different assumptions, formulas and functions different solutions appear.

For easier reading only the male pronouns are used in the whole work, but all examples and arguments apply equally to both genders.

In order to generally explain principal-agent relationships with moral hazard one should think of contract theory which is based on the self-interest hypothesis. So the agent wants to maximize his wealth and other private benefits but this is decreasing in the cost of effort. The principal's aim is to maximize the firm's benefit, but he has to pay the agent his wage. Imagine that in case of delegated bargaining, the firm's wellbeing depends on the agent's action. There the aim of contract theory is the design of incentives so as to induce the agent to behave in a manner the principal wishes. As now also social preferences matter the concern about other's payoffs respectively their wellbeing is important for the design of incentives.¹⁴

¹³ See Itoh, H. (2004), page 18ff

¹⁴ See Itoh, H. (2004), page 19

1.2 Constitution

The rest of the thesis is organized as follows: In Section 2 for better comprehension the nature of social preferences is presented, where the common preference types are discussed. In Section 3 one can see a sample of models considering social preferences as inequity aversion of Fehr et al.¹⁵, Reciprocity of Rabin¹⁶ and so forth. In Section 4 the basis of the models of Lammers¹⁷ and Itoh¹⁸ are described. In Section 5 the two models are subject to a deeper analysis in case of individual work, so the agent works independently. In Section 6 the models are analyzed in case of interaction with a third party. Both sections are concluded with a comparison. In Section 7 finally the work is concluded with a summary.

2 The nature of social preferences

A lot of people possess social preferences, so they are, beside their own material self-interest, concerned about other's material payoff. So it is important to take social preferences into account. This is to understand effects of competition on market outcomes, forces shaping market failures, laws governing cooperation and common action, effects of material incentives and which contracts are optimal, and so forth. The important types of social preferences are described to better understand the deeper analysis in Section 4, 5 and 6: Preference for reciprocity or reciprocal fairness, inequity aversion and altruism are mentioned.¹⁹

¹⁵ See Fehr et al. (1999), page 817ff

¹⁶ See Rabin, M. (1993), page 1281ff

¹⁷ See Lammers, F. (2010), page 169ff

¹⁸ See Itoh, H. (2004), page 18ff

¹⁹ See Fehr, E., Fischbacher, U. (2002), page 1f

2.1 Reciprocity

A person being reciprocal responds kindly to actions being kind and adversely to actions being hostile. The feeling if an action is seen as kind or hostile depends on the fairness of the consequences and the intention associated with the action. This is determined by the fair distribution of the payoffs. But reciprocity is not forced by the expectation of future material benefit, as the reciprocal person is responding to friendly or hostile actions even if he can't expect any material gain. Models dealing with reciprocity have been established for example by Rabin²⁰, Charness and Rabin²¹, Levin²², Falk and Fischbacher²³, Dufwenberg and Kirchsteiner²⁴ and so on.²⁵

2.2 Fairness – Inequity aversion

Fehr and Schmidt²⁶ or Bolton and Ockenfels²⁷ established models including fairness concerns respectively inequity aversion. Persons who are inequity averse want to achieve an equal distribution of material payoffs, so they are altruistic towards other persons. In other words, inequity averse persons want to increase the other's material payoffs if they are below an equitable benchmark. Conversely they feel envy if the income of the other person is above this benchmark. Inequity averse persons and reciprocal persons behave in many situations similarly because they use the same notion of a fair or equitable payoff. As an example, imagine that the fair-minded person wants to reduce the other person's payoff if this person made a decision so that the payoff of the reciprocal or inequity averse person is lower than his payoff. However according to some

²⁰ See Rabin, M. (1993), page 1281ff

²¹ See Charness, G., Rabin, M. (2000), page 1ff

²² See Levin, D.K. (1998), page 593ff

²³ See Falk, A., Fischbacher, U. (2006), page 293ff

²⁴ See Dufwenberg, M., Kirchsteiner, G. (2004), page 268ff

²⁵ See Fehr, E., Fischbacher, U. (2002), page 2f

²⁶ See Fehr, E., Schmidt, K. (1999), page 817ff

²⁷ See Bolton, G.E., Ockenfels, A. (2000), page 166ff

evidence, reciprocity is the stronger and quantitatively more important fairness consideration; nevertheless inequity averse models are simpler than reciprocal models. For that reciprocity is often copied by inequity aversion.²⁸

2.3 Altruism

Different from the previous two types, altruism is a form of unconditional kindness. An altruistic person never chooses an action that decreases the payoff of other persons. In other words altruism doesn't develop because of received altruism. Furthermore people are willing to punish other people for their hostile or unfair actions. As it is a form of unconditional kindness, it can't explain the phenomenon of conditional cooperation, so for example why people are willing to raise their voluntary cooperation in response to cooperation of the other's. There are also research results that people possess spiteful or envious preferences according to the model of Falk et al²⁹. Spiteful or envious people act selfishly and they want to decrease the material payoff of other persons not depending on the fair or unfair behavior of these persons or the payoff distribution. In fact spitefulness is less important than reciprocity. Furthermore altruism and spitefulness can't explain why the same (!) people act in diverse situations differently – either they want to help or harm them. Models including altruism are for example Andreoni³⁰, and some others.³¹

²⁸ See Fehr, E., Fischbacher, U. (2002), page 3

²⁹ See Falk et al. (2000a), page 1ff

³⁰ See Andreoni, J. (1989): page 1447ff

³¹ See Fehr, E., Fischbacher, U. (2002), page 3f

3 Overview of models considering social preferences

Fehr/Fischbacher/Schmidt: Inequity Aversion

In this case fairness is modeled as self-centered inequity aversion. Inequity aversion means that individuals don't like inequitable outcomes and they become self-centered if they care about the fairness of their own material payoff in relation to the other's payoff. Furthermore the preference type is determined by the economic environment.³² In the model of Fehr et al. (1999) the individuals can induce, under specific competitive conditions, other actors being extremely inequity-averse playing selfishly. This is a contradiction to the model of Lammers where the preference type is seen as stable which will be seen later in the model description³³. Many experiments have proofed that individuals resist outcomes which are either to their disadvantage or to their advantage.

In Figure 1 one can see the individual's utility as a function of his payoff.³⁴

³² See Fehr, E., Fischbacher, U. (2002), page 3

³³ See Lammers, F. (2010), page 170

³⁴ See Fehr, E. Fischbacher, U., Schmidt, K.M. (1999), page 819

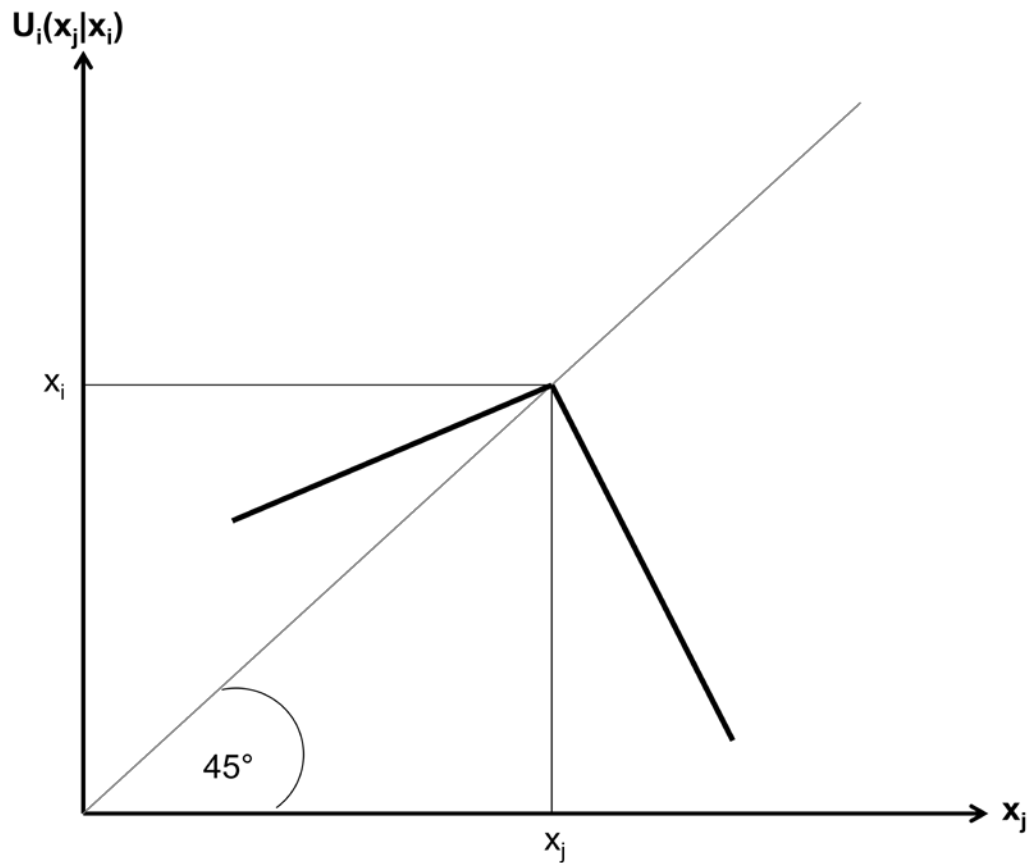


Figure 1: Utility function Inequity Aversion³⁵

The maximum utility is reached at $x_j = x_i$. As one can see, the disutility of being behind (agent's income is lower than that of the others) is bigger than the disutility of being ahead (agent's payoff is higher than the others income).³⁶

Their model is fundamental for many other models concerning social preferences, as for example Itoh³⁷, Neilson and Stowe³⁸, Bolton and Ockenfels³⁹ and many others.

³⁵ See Fehr et al. (1999), page 823

³⁶ See Fehr, E. Fischbacher, U., Schmidt, K.M. (1999), page 823

³⁷ See Itoh, H. (2004), page 18ff

³⁸ See Neilson, W.S., Stowe, J. (2003), 1ff

Bolton/Ockenfels: Inequity Aversion

In the model of Bolton and Ockenfels⁴⁰ similar predictions to the previous mentioned model of Fehr and Schmidt can be found⁴¹. Actors want an equal distribution among all players and do not only consider their own relative payoff. In this model, the actors dislike payoffs which aren't close to the average payoffs; this is the main difference to the model above.⁴²

Rabin: Reciprocity

“According to this theory, reciprocity is a behavioral response to perceived kindness and unkindness, where kindness comprises both distributional fairness as well as fairness intentions.”⁴³

Describing this citation, reciprocity means that people act altruistically to people being even altruistically and want to hurt those who hurt other people. Furthermore, the smaller the sacrificing cost, the greater the effect of these two motivations. The payoff depends on the one hand on the agent's belief and on the other hand on his action. Of course considering the belief complicates the model. So the agent's action is kind or unkind depending on the consequences he thinks of his action. So the agent's utility depends on his strategy, his belief

³⁹ See Bolton, G.E., Ockenfels, A. (2000), page 166ff

⁴⁰ See Bolton, G.E., Ockenfels, A. (2000), page 166ff

⁴¹ See Fehr, E., Schmidt, K. (1999), page 817ff

⁴² See Bolton, G.E., Ockenfels, A. (2000), page 166ff

⁴³ Falk, A., Fischbacher, U. (2006), page 294

which strategy the other players choose and his belief about what the others believe his strategy will be.⁴⁴

Nevertheless, “in many situations reciprocal persons and inequity averse persons behave in similar ways”.⁴⁵ The reason is because both models use the same notion of a fair respectively equitable payoff. Again, models concerning inequity aversion are seen as simpler, thus models of reciprocity often imitate inequity averse behavior.⁴⁶

Mayer/Pfeiffer: Principles of incentives in case of risk aversion and social preferences

In this model the standard LEN model is enhanced by social preferences:

L	Linear technology, contract and social preference
E	Exponential utility
N	Normal distribution
S	Social preference

The agent compares ex post his income with the principal's payoff and feels utility respectively disutility because of his social preferences (envy or spitefulness). So a preference coefficient is included in the utility function which measures the degree of envy respectively spitefulness. In conclusion, if the preference coefficient is high, the agent wants an equal distribution of the principal's and his payoff. Furthermore, if the agent feels envy, the principal has to pay him a higher wage, but he can also benefit by inducing the agent to exert more effort. In contrast, if the agent feels spitefulness, so he expects to be “ahead”, the principal can again benefit by inducing the agent to choose the high effort.⁴⁷

⁴⁴ See Rabin, M. (1993), page 1281ff

⁴⁵ Fehr, E., Fischbacher, U. (2002), page 3

⁴⁶ See Fehr, E., Fischbacher, U. (2002), page 3

⁴⁷ See Mayer, B., Pfeiffer, T. (2004), page 1047ff

In the following a deeper analysis of the two models is made. The purpose is the expectation of finding out which agent will be preferred – fair or selfish. The next section starts with the description of both models as an introduction.

4 Description of the models

4.1.1 Introduction

The first model is based on the article „Fairness in Delegated Bargaining“ (Spring 2010) of Frauke Lammers published in Journal of Economics & Management Strategy, Volume 19 and on her working paper also called “Fairness in Delegated Bargaining” (2007) not published officially. The other model is based on the article „Moral hazard and other-regarding preferences“ (March 2004) of Hideshi Itoh published in the Japanese Economic Review, Volume 55, Number 1.

For better comparability in the whole thesis the denoted variables from the original model of Itoh were replaced by the variables of the model of Lammers.

The model of Lammers models social preferences of agents and tries to find out whether the principal chooses to hire a fair (altruistic) or selfish one. This answer seems very straightforward, as he needs fewer incentives to act in the way the principal wants to.⁴⁸

The model of Itoh is based on the model “A theory of fairness, competition, and cooperation”⁴⁹, respectively “ERC: A theory of equity, reciprocity and competition”⁵⁰ as they present an alternative formulation of inequity aversion. “Their model is an example of the distributional approach to other-regarding preferences in which only the final monetary distribution among players matter”⁵¹. So the agent feels guilty if his income is above that of the other agents, but he

⁴⁸ See Lammers, F. (2010), page 169

⁴⁹ See Fehr, et al. (1999), page 817ff

⁵⁰ See Bolton, G.E., Ockenfels, A. (2000), page 166ff

⁵¹ Itoh, H. (2004), page 20

feels also envious if the other agent's payoff is above his payoff. But Itoh⁵² extends this basic model by competitive or status-seeking preferences, so that the self-interested agent loves being ahead (his income is above the others), but the fair agent disrelishes being behind (his income is below the others).⁵³

The two models are compared because they have a lot of similarities. In both models the agent has to spend effort to produce the required output and he can decide whether to choose the high or low effort. Of course the principal prefers in both models the high effort as it increases his payoff. But both models, using the utility function as a decision factor, result in different solutions. As this is a very interesting point the models are analyzed according to their decision factors including distribution of wage payments, equal timing, and extent of care for third party and especially for principal and so forth.

4.1.2 The models^{54 55}

As can be seen in Figure 2 and based on the classical principal-agent-model, the risk-neutral principal "P" hires a risk-neutral agent "A" to work for him either individually or by interacting with a risk-neutral third party "B" (e.g. a buyer or a colleague) on some project.

⁵² See Itoh, H. (2004), page 18ff

⁵³ See Itoh. H. (2004), page 20

⁵⁴ See Itoh. H. (2004), page 18ff

⁵⁵ See Lammers, F. (2007), page 169ff

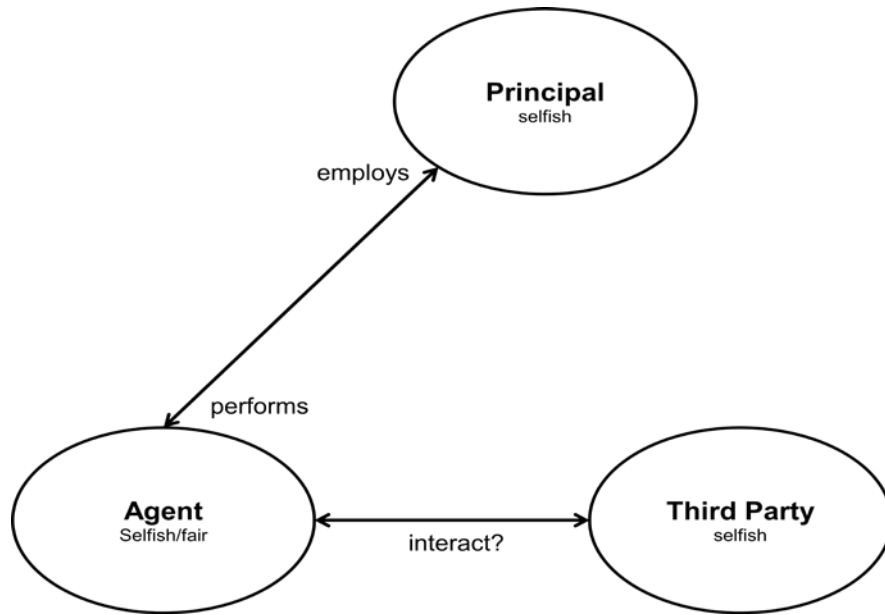


Figure 2: Principal-Agent-Model with fairness considerations⁵⁶

This interaction results in expected monetary payoffs to all involved parties:

- π_P (principal's payoff),
- π_A (agent's payoff), and
- π_B (third party's payoff).

Generally output is stochastic and can take the values $\theta \in \{\theta_L, \theta_H\}$ with $\theta_L = 0$ or $\theta_H = 1$. With effort $e \in [0, 1]$, which is unobservable to the principal, the agent increases the principal's expected payoff. The cost of effort is $c(e) = \frac{c}{2}e^2$ with $c > 0$. For simplicity cost of effort is denoted as cost c .

⁵⁶ Own illustration

Setting	Output	$q'_A(e)$	$q'_B(e)$	P's access to
Individual work	$q = q_A$	$q'_A(e) > 0$	$q'_B(e) = 0$	q_A
Internal third party	$q = q_A + q_B$	$q'_A(e) > 0$	$q'_B(e) < 0$	$q_A + q_B$

Table 1: The settings⁵⁷

In table 1 the settings, individual work and interaction with an internal third party, are shortly presented. The output variable can be differently interpreted: In the case of individual work, it represents the agent's production q_A accruing to the principal. The agent's effort level is equal the probability that a high outcome θ_H is realized ($q_A = e\theta_H$) and the higher his effort, the higher the outcome. In contrast to the case when the agent interacts with a third party, output represents the third party's valuation and the agent's effort can have an impact on the price he is able to negotiate. In this case the two agents generate output separately and the higher the agent's effort, the higher is the outcome but the agent lowers the third party's output.

According to social preferences, it is assumed that the principal and the third party are selfish. In contrast, the agent can either be selfish or fair.

The utility of the principal is

$$u_P = q_j - w_j$$

The selfish person cares only for his own well-being. In this case his utility is

$$u_S = U_S(\pi_A) \rightarrow \\ U_S = \pi_A$$

⁵⁷ See Lammers, F. (2007), page 8

A fair agent, additionally to his own monetary payoffs, cares to some extent for the payoffs of the principal and the third party, which acts as an add-on to his utility.

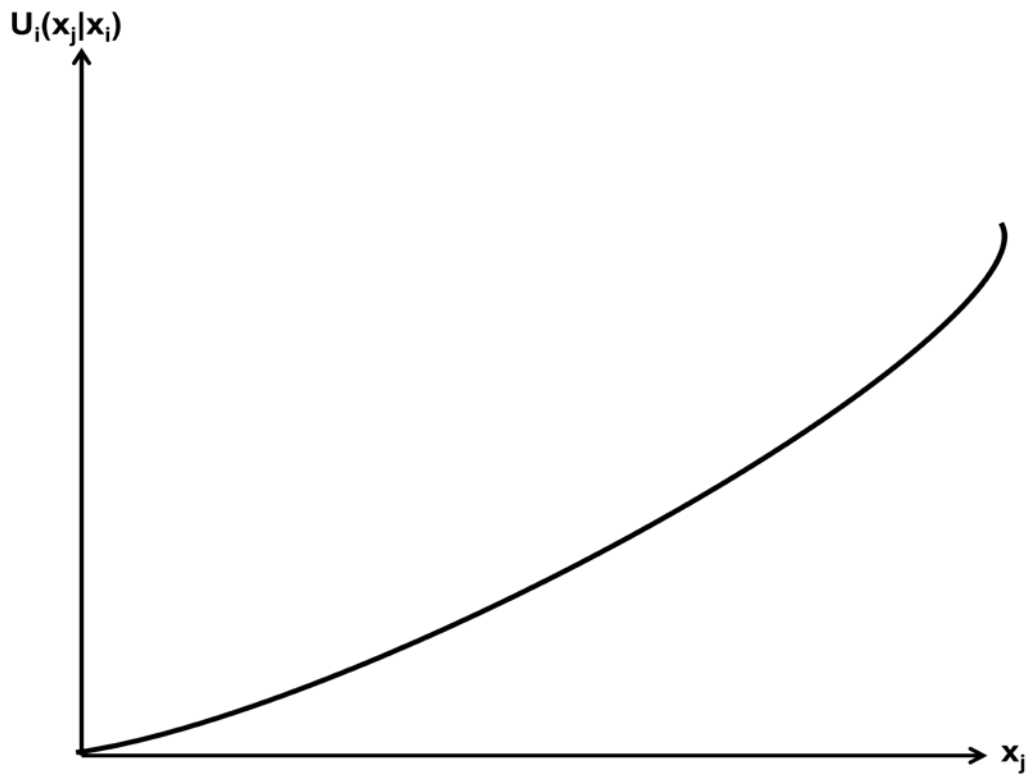


Figure 3: Convex utility function

As can be seen from Figure 3, in the model of Lammers the agent's convex utility function is then

$$u_F = U_F(\pi_A, \pi_P, \pi_B)$$

The agent cares with the factors $\mu_P, \mu_B > 0$ for the principal's payoff (μ_P) and for the third party's payoff (μ_B). Therefore his utility is stated as

$$U_F = \pi_A + \mu_P \pi_P + \mu_B \pi_B$$

With the assumption

$$\mu_P, \mu_B < \frac{1}{2}$$

it is ruled out that the agent needs no monetary compensation for his work as he would receive such a high utility from the principal's and/or buyer's payoff. But this is rather the case for voluntary activities and not relevant in this thesis.

Furthermore the factors μ_P and μ_B can (not have to!) be positively correlated, as an agent who has a high concern for the principal might also have a high concern for the third party. An agent, who feels like an insider, doesn't require much monetary compensation to choose a high effort level⁵⁸. For $\mu_P \gg \mu_B$ the agent can be interpreted as being an insider as he cares more about the principal's payoff than the third party's payoff.

A strong assumption is, that the principal knows the type of the agent and the values of μ_P and μ_B . Nevertheless, many companies learn about the agent's type by conducting personality tests like the Big 5⁵⁹ or the Myer-Briggs Type Indicator⁶⁰.

⁵⁸ Akerlof, G.A., Kranton, R.E. (2005), page 11

⁵⁹ See Simon, W. (2006), page 113ff

⁶⁰ See Simon, W. (2006), page 299ff

The timing of the game is illustrated in Figure 4 and the following:

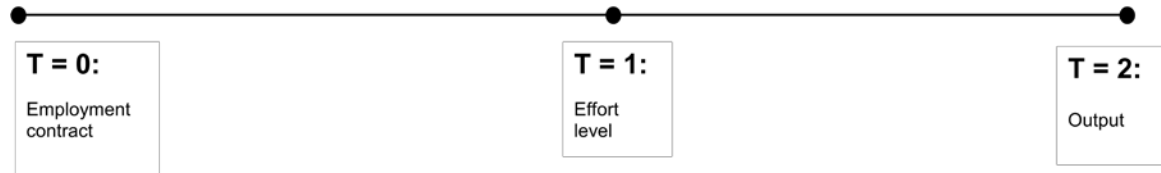


Figure 4: Timing⁶¹

$T = 0$: The agent gets offered an employment contract of the principal and then he can accept or reject it. If the agent accepts the contract, the sequence will be continued.

$T = 1$: The agent chooses his effort level.

$T = 2$: The output(s) is (are) realized and the payoffs are distributed.

Furthermore the agent is restricted by limited liability and his reservation utility and initial wealth equal 0 (in the sense of no effort, no money). Such contracts with limited liability, as bankruptcy clauses, some conditions under which violation of a contract is allowable and so forth, are common in practice⁶².

As there are only two output-levels and for simplicity a linear employment contract is assumed including a fixed component and a variable component depending on the agent's output:

$$w_j = \alpha + \beta\theta \text{ where } j = s, f$$

⁶¹ Own illustration

⁶² Sappington, D., (1983), page 2

In the model of Itoh, if the project succeeds with probability p_1 it generates high output θ_H , otherwise if it fails with probability $p_0 = 1 - p_1$ it generates low output θ_L and $1 > p_1 > p_0 > 0$. The agent receives wage w_s if the project succeeds and w_f if it fails

$$w_j \geq 0 \text{ where } j = s, f$$

Furthermore the incentive compatibility constraint and participation constraint have to be considered:

The incentive compatibility is

$$(w_s - w_f)(p_1 - p_0) \geq c \rightarrow$$

$$\Delta_p \Delta_w \geq c \text{ (IC)}$$

and the participation constraint is denoted by

$$w_f + p_1(w_s - w_f) \geq c \rightarrow$$

$$w_f + p_1 \Delta_w \geq c \text{ (PC)}$$

In fact, the principal selects the wage $(w_s, w_f) \in \mathcal{C}$ that minimizes his expected payment $w_f + p_1 \Delta_w$ or increases his payoff.

In the model of Itoh, the fair agent's utility function is the following and illustrated in Figure 1:

$$u_A = w_j - c_i - \mu_P v(\max\{q_j - 2w_j, 0\}) - \mu_P \gamma v(\max\{2w_j - q_j, 0\})$$

$\mu_P \geq 0$ and γ are constants and it is assumed that $v(0) = 0$, $v'(z) > 0$ for all $z > 0$ and $\lim_{z \rightarrow \infty} v(z) = \infty$.

5 Individual work without interaction

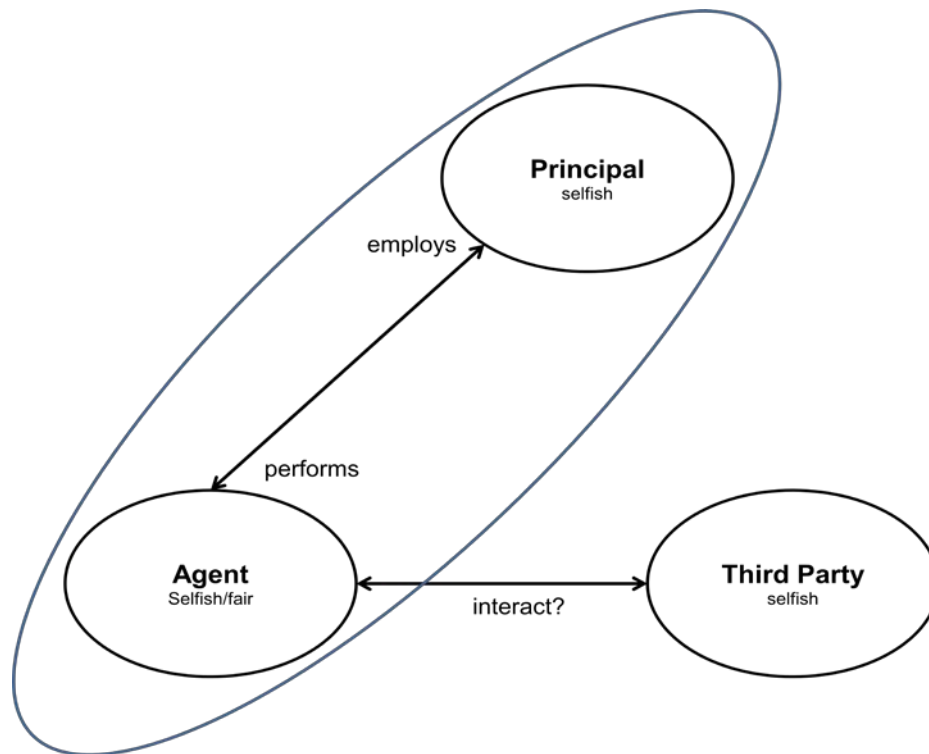


Figure 5: Individual Work without Interaction⁶³

As already mentioned, different situations will be analyzed, starting with no interaction between more agents. In the following topic only one agent works for the principal without interacting with another agent which is illustrated in Figure 5.

⁶³ Own illustration

5.1 Fairness in delegated bargaining⁶⁴

This topic is about individual work when no third party is participating. The agent exerts effort $e \in [0, 1]$ to produce output $q = q_A = e\theta_H$. The agent can either be selfish or fair, where the fair agent is additionally to his own monetary payoff concerned about the principal's payoff.

The principal's payoff is

$$\begin{aligned}\pi_P &= q_A - E(w) \rightarrow \\ \pi_P &= e\theta_H - E(w)\end{aligned}$$

and the agent's is

$$\begin{aligned}\pi_A &= E(w) - c(e) \rightarrow \\ \pi_A &= E(w) - \frac{c}{2}e^2\end{aligned}$$

where $E(w)$ represents the expected wage payment.

5.1.1 Choice of effort

Now it has to be checked whether the principal always prefers to hire a fair agent in case of individual work. By backward induction, the optimal employment contract can be found. Given the agent's utility functions the optimal effort levels are:

$$\begin{aligned}U_S = \pi_A = E(w) - c(e) &\rightarrow \pi_A = \alpha + \beta q_A - \frac{c}{2}e^{NI}(S)^2 \rightarrow \\ \pi_A &= \alpha + \beta e^{NI}(S)\theta_H - \frac{c}{2}e^{NI}(S)^2 \rightarrow \\ e^{NI}(S) &= \beta \frac{\theta_H}{c}\end{aligned}$$

⁶⁴ See Lammers, F. (2010), page 173ff

for the selfish agent and

$$\begin{aligned}
U_F &= \pi_A + \mu_P \pi_P \rightarrow \\
&E(w) - c(e) + \mu_P [q_A - E(w)] \\
&\rightarrow \left(\alpha + \beta e \theta_H - \frac{c}{2} e^{NI}(F)^2 \right) + \mu_P [e \theta_H - (\alpha + \beta e \theta_H)] \rightarrow \\
e^{NI}(F) &= (\beta + \mu_P (1 - \beta)) \frac{\theta_H}{c}
\end{aligned}$$

for the fair agent.

As one can see, because of the fairness concern $\mu_P(1 - \beta) > 0$, therefore the fair agent initiates more effort than the selfish one, thus

$$e^{NI}(F) > e^{NI}(S)$$

Furthermore his effort rises the more he cares for the principal (the higher μ_P , the higher $e^{NI}(F)$).

Nevertheless there are 2 cases identified where both agents exert the same effort level, $e^{NI}(F) = e^{NI}(S)$. First, if $\beta = 1$; unfortunately in this case the agent would be the owner of the firm and his concern for the principal's payoff seems irrelevant. The other case, if $\mu_P = 0$, so the agent feels no concern for the principal, therefore he equals the selfish agent.

5.1.2 Optimal employment contracts

Given the limited liability constraint of the agent, the principal formulates the optimal employment contract for the selfish agent by maximizing his payoff:

$$\begin{aligned}
\max \pi_P &= q_A - E(w) \rightarrow e \theta_H - (\alpha + \beta e \theta_H) \rightarrow \\
&(1 - \beta) e^{NI}(S) \theta_H - \alpha \rightarrow \\
&\beta (1 - \beta) \frac{\theta_H^2}{c} - \alpha
\end{aligned}$$

$$\text{s.t. } U_S = \pi_A = \alpha + \beta e^{NI}(S)\theta_H - \frac{c}{2}e^{NI}(S)^2 \rightarrow \alpha + \frac{(\beta\theta_H)^2}{2c} \geq 0 \text{ and } \alpha \geq 0$$

And for the fair agent:

$$\max \pi_P = (1 - \beta)e^{NI}(F)\theta_H - \alpha \rightarrow$$

$$[\beta(1 - \beta) + \mu_P(1 - \beta)^2] \frac{\theta_H^2}{c} - \alpha$$

$$\text{s.t. } U_F = \alpha + \beta e^{NI}(F)\theta_H - \frac{c}{2}e^{NI}(F)^2 + \left(\mu_P((1 - \beta)e^{NI}(F)\theta_H - \alpha)\right) \rightarrow$$

$$\alpha(1 - \mu_P) - \left(\beta + \mu_P(1 - \beta)\right)^2 \frac{\theta_H^2}{2c} \geq 0 \text{ and } \alpha \geq 0$$

After using the Lagrange-function, one can see that the principal offers a contract with no fixed payment to both agents

$$\alpha^{NI} = 0$$

and a variable payment of

$$\beta^{NI}(S) = \frac{1}{2}$$

to the selfish agent and the fair agent receives a variable component of

$$\beta^{NI}(F) = \frac{1 - 2\mu_P}{2(1 - \mu_P)}$$

As the fair agent gets extra utility of $\mu_P\pi_P$ of the interaction, he needs a lower commission rate than the selfish agent ($\beta^{NI}(S) > \beta^{NI}(F)$). As can be seen, for $\mu_P = \frac{1}{2}$ the fair agent needs no monetary incentives ($\beta^{NI}(F) = 0$). For higher values of μ_P the agent would also be willing to pay the principle for working. This is in this setting irrelevant because the agent is restricted by limited liability and in turn the assumption made above ($\mu_P, \mu_B < \frac{1}{2}$) holds.

Summing up, the principal's expected payoff increases the more the agent cares for him and is

$$\pi_P^{NI}(S) = \frac{\theta_H^2}{4c}$$

for the selfish agent and

$$\pi_P^{NI}(F) = \frac{\theta_H^2}{4c(1 - \mu_P)}$$

for the fair agent.

5.1.3 Choice of agent

Of course the principal's payoff with a fair agent is higher,

$$\pi_P^{NI}(F) > \pi_P^{NI}(S),$$

so the fair agent will be preferred as for individual work the fair agent needs less monetary incentives as he gets extra utility from interaction and furthermore the selfish agent has no advantages over the fair one.

5.2 Moral Hazard and other-regarding preferences⁶⁵

5.2.1 Self-interested agent

The optimal contract (w_s^*, w_f^*) is unique as it minimizes the principal's expected payment to the agent and is given from the incentive compatibility constraint by

$$w_f^* = 0$$

because of limited liability of the agent and

⁶⁵ See Itoh. H. (2004), page 25ff

$$\Delta_p \Delta_w = c \rightarrow \Delta_p (w_s^* - w_f^*) = c \rightarrow \Delta_p (w_s^* - 0) = c \rightarrow \Delta_p w_s^* = c \rightarrow$$

$$w_s^* = \frac{c}{\Delta_p}$$

The principal's expected payment is therefore

$$w_f + p_1 \Delta_w \rightarrow 0 + p_1 w_s^* \rightarrow$$

$$p_1 w_s^*$$

One sees that the agent receives only a wage if the project succeeds and in fact the agent has an impact on his wage payment by e .

In order to show the uniqueness it is supposed that another contract $(w_s, w_f) \neq (w_s^*, 0)$ solves the principal's problem. As the incentive compatibility constraint must bind at the optimum,

$$\Delta_w = w_s - w_f = w_s^*$$

and

$$w_f > 0$$

holds. Then the principal's expected payment equals

$$w_f + p_1 \Delta_w \rightarrow$$

$$w_f + p_1 w_s^* > p_1 w_s^*$$

which is inconsistent with the optimality of (w_s, w_f) as the principal has to pay the agent more.

So the principal has to provide the self-interested agent high-powered incentives, so as to act in the interest of the principal, and leave him rents, because of the moral hazard problem:

$$u_A = p_1 w_s^* - c \rightarrow p_1 \frac{c}{\Delta_p} - c \rightarrow$$

$$p_0 \frac{c}{\Delta_p} > 0$$

5.2.2 Other-regarding agent

In this case the agent compares his income w_j with the principal's payoff $q_j - w_j$ and for the time being doesn't take into account the disutility of his action c_i .

The utility function can be rewritten to:

$$u_A = \begin{cases} w_j - c_i - \mu_P \gamma v(2w_j - q_j) & \text{if } w_j \geq q_j - w_j \\ w_j - c_i - \mu_P v(q_j - 2w_j) & \text{if } w_j \leq q_j - w_j \end{cases}$$

If the last term $2w_j - q_j > 0$ then the principal's monetary payoff is smaller than the agent's one, where the agent is "ahead". In opposite, the agent is "behind" if $q_j - 2w_j > 0$.

The constant μ_P measures the extent to which the agent cares about the principal's material payoff. If $\gamma < 0$ the agent wants to increase the difference in payoffs when he is "ahead", this means he is competitive or status seeking (= selfish). In opposite if $\gamma > 0$ the agent seems to be inequity averse and his utility is decreasing the higher the difference in payoffs between the agent and the principal, whether he is ahead or behind.

Based on some experimental evidence $\gamma \leq 1$, as the agent suffers more from inequality when he is behind than when he is ahead.⁶⁶ But at this point this restriction is not considered in that way, because it can be seen how the extent affects the incentives with the factor μ_P .

⁶⁶ See Fehr, E., Schmidt, K. (1999), page 822

5.2.3 Wage if project fails is 0

For better comparison of the next section, this section is an introduction. For simplicity actually $w_f = 0$ is assumed, because of imposing the limited liability constraint ($w_s \leq q$, $w_f \leq 0$). Afterwards it is shown that the contract $w_f = 0$ in fact is optimal if an additional assumption instead of the limited liability constraint is introduced.

As the agent in this case doesn't suffer from inequality when the project fails (as no one receives any payment), his other-regarding (fairness) preferences deal with how the benefit of the project is split between the principal and the agent. The incentive compatibility constraint can thus be rewritten as

$$u_A = \begin{cases} w_s - \mu_P \gamma v(2w_s - q_s) & \text{if } w_s \geq \frac{q}{2} \\ w_s - \mu_P v(q_s - 2w_s) & \text{if } w_s \leq \frac{q}{2} \end{cases} \geq \frac{c}{\Delta p} = w_s^* \quad \begin{array}{l} \text{IC2a} \\ \text{IC2b} \end{array}$$

where (IC2a) is the case when the agent is ahead and (IC2b) represents the case when he is behind. Furthermore to act fairly the utility has to be greater than or equal to the self-interested agent's utility.

To check whether a fair or selfish agent is preferred first a necessary and sufficient condition is assumed that there exists a contract which satisfies the case when the agent is behind (IC2b, selfish) if

$$\frac{1}{2} \Delta_p q \geq c$$

To proof this one can see from the incentive compatibility constraint (IC2b) that the left hand side of the equation in case of being behind increases with w_s and the existence is guaranteed if $w_s = \frac{q}{2}$, that is if $\frac{1}{2} \Delta_p q \geq c$ holds and because $v(0) = 0$ ($\frac{c}{\Delta p} \leq \frac{q}{2}$). In opposite, if $w_s \leq \frac{q}{2}$ satisfies the incentive compatibility constraint then the following inequality holds:

$$w_s^* = \frac{c}{\Delta p} \leq w_s - \mu_P v(q - 2w_s) \leq \frac{q}{2}$$

which contradicts the incentive compatibility constraint as the agent's utility is then lower as in the case of self-interest.

First $\hat{w}_s$ has to be defined by the adjusted incentive compatibility constraint

$$\hat{w}_s - \mu_P v(q_s - 2\hat{w}_s) = \frac{c}{\Delta p}$$

If $\frac{1}{2}\Delta_p q \geq c$ holds, $\hat{w}_s$ satisfies

$$\hat{w}_s \leq \frac{q}{2}$$

Since the principal focuses on the high effort e_1 , we see the inequality $\frac{1}{2}\Delta_p q \geq c$ as an assumption.

The principal's profit of the successful project is large enough that $\frac{1}{2}\Delta_p q \geq c$ holds. (**Assumption 1**)

Then (i) the first-best contract $(\hat{w}_s, 0)$ is optimal as the high effort is enforceable and (ii) the principal's expected payment to the agent is increasing in the agent's concern for him (μ_P). In other words the more the agent cares for the principal, the more he wants to be compensated for his "suffering".

(i) To proof this one shows that this satisfies the participation constraint:

$$w_s - \mu_P v(q - 2w_s) \geq \frac{c}{p_1} = \bar{w}_s \text{ (PCb)}$$

Since by definition

$$\hat{w}_s - \mu_P v(q_s - 2w_s) = \frac{c}{\Delta p} > \frac{c}{p_1}$$

the participation constraint is thus satisfied as the principal can enforce the agent to choose the high effort.

- (ii) As the principal's payment is $p_1 \hat{w}_s$ it is adequate to show that $\hat{w}_s$ is increasing in μ_P by defining the left hand-side of $\hat{w}_s - \mu_P v(q_s - 2\hat{w}_s) = \frac{c}{\Delta p}$ by derivation. Then

$$\frac{\partial f(\mu_P)}{\partial \mu_P} = -v(q_s - 2\hat{w}_s) < 0$$

for $\hat{w}_s < \frac{q}{2}$ which states that it is strictly increasing in μ_P . Otherwise if $\frac{1}{2} \Delta_p q = c$ holds it doesn't depend on μ_P .

In fact it can be seen that,

- (i) as in the case of pure self-interest, at optimum the incentive compatibility constraint binds, while the participation constraint does not and thus the selfish agent earns rents as he only cares about his wage and only participates if the payoff exceeds the reservation utility.
- (ii) Furthermore the principal is worse off the more the agent cares about his wellbeing, as the left-hand side of the incentive compatibility constraint in case of being behind is decreasing in μ_P . The agent suffers from inequity and the principal has therefore to pay more to satisfy the incentive compatibility constraint the more the agent cares for him. In other words, if the agent's income is below the principal's payoff, the principal is worse off, the more other-regarding the agent is. The principal has therefore to pay more the more inequity averse the agent is in order to satisfy the incentive compatibility constraint.

5.2.4 Wage if project succeeds is higher than 0

The following problem analyses the case if contracts with $w_f > 0$ are feasible; then also $(\hat{w}_s, 0)$ is again unique optimal under an additional assumption, means if

$$2\mu_P\gamma v'(z) \leq 1$$

is satisfied for all $z > 0$.

To proof this assume that the previous result doesn't hold and (w_s, w_f) with $w_s, w_f > 0$ is optimal. Therefore the participation constraint

$$w_f + p_1\Delta_w \leq p_1\hat{w}_s$$

holds.

Because $\hat{w}_s \leq \frac{q}{2}$ and $w_f > 0$, $w_s < \frac{q}{2}$ must also hold and (w_s, w_f) satisfies the incentive compatibility constraint

$$\Delta_w - \mu_P v(q - 2w_s) + \mu_P \gamma v(2w_f) \geq \frac{c}{\Delta_p}$$

5.2.5 Choice of agent

Combining the two derived incentive compatibility constraints

$$\Delta_w - \mu_P v(q - 2w_s) + \mu_P \gamma v(2w_f) \geq \frac{c}{\Delta_p} \text{ and}$$

$$\hat{w}_s - \mu_P v(q_s - 2w_s) = \frac{c}{\Delta_p} \text{ yield}$$

$$\begin{aligned} \Delta_w + \mu_P \gamma v(2w_f) &\geq \hat{w}_s + \mu_P v(q - 2w_s) - \mu_P v(q - 2\hat{w}_s) \\ &\geq \Delta_w + \frac{w_f}{p_1} + \mu_P v(q - 2w_s) - \mu_P v(q - 2\hat{w}_s) \end{aligned}$$

Rearranging the second equality with

$$w_f + p_1 \Delta_w \leq p_1 \hat{w}_s \rightarrow$$

$$\hat{w}_s \geq \frac{w_f}{p_1} + \Delta_w$$

one gets

$$-\frac{w_f}{p_1} + \mu_P \gamma v(2w_f) \geq \mu_P v(q - 2w_s) - \mu_P v(q - 2\hat{w}_s)$$

The right hand side of this equation is positive since $w_f + p_1 \Delta_w \leq p_1 \hat{w}_s$ leads to $w_s \leq \hat{w}_s$. So if the left-hand side is coherent, (w_s, w_f) doesn't satisfy the incentive compatibility constraint, which is seen as a contradiction.

In this setting, where $w_f > 0$, a contract even creates inequity if the project fails. Since the agent receives the lower wage, and the principal pays the lower wage w_f , the agent is ahead. If he is now inequity averse ($\gamma > 0$) this increases his incentive to choose μ_P as can be seen from the incentive compatibility constraint so as to be behind. However, to benefit the principal, he must decrease the higher wage w_s so as to lower Δ_w and to raise inequity $q - 2w$ when the agent is behind. So the selfish agent is preferred which brings negative effects on incentives.

If $2\mu_P \gamma v'(z) \leq 1$ is violated, then $2\mu_P \gamma v'(z) > 1$. So the agent is willing to transfer some payment back to the principal so as to be less ahead. For that reason $1 \geq 2\mu_P \gamma v'(z)$ is assumed. The condition holds even if the agent is status-seeking ($\gamma < 0$).

When Assumption 1 does not hold

Instead now it is assumed that

$$\frac{1}{2}\Delta_p q < c$$

As it is known from the previous sector, there is no contract $(w_s, 0)$ that satisfies the incentive compatibility constraint in the range of $w_s \leq \frac{q}{2}$. For that the principal now has to choose $w_s > \frac{q}{2}$. Again the left-hand side of (IC2a) is increasing in w_s and w_s^+ is defined by

$$w_s^+ - \mu_p \gamma v(2w_s^+ - q_s) = \frac{c}{\Delta p}$$

If this condition is now differentiated with respect to μ_p one receives

$$\frac{\partial f(\mu_p)}{\partial \mu_p} = -\gamma v(2w_s^+ - q)$$

for $w_s^+ > \frac{q}{2}$.

So w_s^+ is strictly increasing in μ_p if $\gamma > 0$ - so the agent is inequity averse - and strictly decreasing if $\gamma < 0$ - so the agent is status-seeking. Again, the principal is worse off the more the fair-minded agent cares for him if he is inequity averse. Conversely, if the agent is selfish, the principal can benefit more as with an agent being other-regarding. The status-seeking agent likes being ahead and so the principal can implement the effort e_1 at lower costs the more competitive the agent seems.

5.2.6 “Net” Material Payoff

In this section the agent compares his “net” material payoff $w_j - c_i$ to the principal's payoff $q_j - w_j$. Reason for that is that the agent knows that he has to

invest cost c_i , so it can be that he is concerned about inequity differently, depending on his chosen effort e_0 or e_1 .

First it is supposed that $q - w_s > w_s$ (so the agent is behind) and for simplicity $w_f = 0$. The incentive compatibility constraint and the participation constraints are therefore

$$\Delta_p w_s - [p_1 \mu_P v(q - 2w_s + c) - p_0 \mu_P v(q - 2w_s)] - (1 - p_1) \mu_P v(c) \geq c \text{ (IC3b)}$$

$$p_1 w_s - p_1 \mu_P v(q - 2w_s + c) - (1 - p_1) \mu_P v(c) \geq c \text{ (PC3b)}$$

The cost of action c is only included in the inequity aversion term in the case of successful project $\mu_P v(q - 2w_s + c)$. If the agent chooses e_0 , the cost is clearly zero and so the change in the inequity aversion term is

$$p_1 \mu_P v(q - 2w_s + c) - p_0 \mu_P v(q - 2w_s)$$

This change is now larger than if Assumption 1 holds. So he can save these costs and therefore minimize the disutility from inequity when he is behind. Differently said, the agent feels less envious if he is behind by choosing e_0 as if he chooses e_1 .

Furthermore a new inequity-averse term $\mu_P v(c)$ when the project fails is introduced. Since the agent has to bear the cost c , he is behind and suffers from inequity. So by choosing e_0 he can again facilitate what he would suffer if the project fails.

So the incentive compatibility constraint and the participation constraints are in fact harder to satisfy and the principal is therefore again worse off the more other-regarding (fair) the agent is.

If it is now supposed that the agent is ahead as the principal's payoff is below his income $q - w_s < w_s - c$ the case with the inequity averse term $\mu_P v(c)$ if the project fails remains, so choosing e_0 relieves the agent of the envy he would

suffer as he would choose the high effort. The incentive compatibility constraint's term changes to

$$p_1\mu_P\gamma v(2w_s - q - c) - p_0\mu_P\gamma v(2w_s - q)$$

where it can be seen that the inequity constant γ is included. If the agent now chooses e_1 he can reduce the difference by c when he is ahead. If the agent is status-seeking ($\gamma < 0$) both constraints are harder to satisfy. If the agent is inequity averse ($\gamma > 0$), it benefits the principal as the incentive compatibility is easier to satisfy, if the other change ($\mu_P v(c)$) doesn't dominate.

5.3 Comparison of "No Interaction"

"Fairness in Delegated Bargaining"	"Moral Hazard and other-regarding preferences"
Fair agent	Basically selfish agent, except if cost of effort is included: rather fair agent (but not always)

Table 2: Comparison of "No Interaction"

In Table 2 the cases of no interaction are compared. According to the model of Lammers⁶⁷ the effort levels of the agents depend on their utility function which in turn depends on the wage w , including a fixed α and variable component β .

Because of the fairness concerns the fair agent spends more effort and his effort raises the more his concern for the principal increases. This is due to the convex

⁶⁷ See Lammers, F. (2007), page 1ff and Lammers, F. (2010), page 169ff

utility function $U_F = \pi_A + \mu_P \pi_P + \mu_B \pi_B$. Both agents (selfish and fair) receive no fixed payments and different variable payments as it depends on their fairness concerns. As the fair agent feels extra utility of his fairness concern for the principal of the interaction, the fair agents needs a lower commission rate than the selfish agent which receives $\frac{1}{2}$ of the output reached. Namely, the fairness concern acts as an add-on to his utility. So the principal has to pay the fair agent less which in turn increases his payoffs.⁶⁸

According to the model of Itoh⁶⁹ the agent compares his payoff (not including his cost c) with the principal's income and this turns out whether the agent is "ahead" (selfish) or "behind" (fair). The inequity averse has to feel a higher utility than the selfish agent in order to act in principal's way. The case when the agent is behind exists if he receives $\frac{1}{2}$ of the output generated which equals the model of Lammers⁷⁰ and this has to be higher than the agent's cost. The principal has again to pay the self-interested agent rents and leave him therefore high-powered incentives. Furthermore, in the single-agent case the principal generally does not benefit from agent's other-regarding preferences due to the utility function concerning inequity aversion.

If the output of the successful project is large, the principal clearly prefers a status-seeking, self-interested agent, which differs from the model of Lammers⁷¹. This is due to the agent's utility function. It is more costly for the principal to induce the other-regarding agent to exert more effort. Compared to the model of Lammers⁷², where the agent feels additional utility so as to exert more effort the inequity averse agent in the model of Itoh reaches his optimum utility at $x_j = x_i$ - so there is a break-even point of utility. Trespassing this benchmark reduces the agent's utility that's why the agent customizes his effort level.

⁶⁸ See Lammers, F. (2010), page 169ff and Lammers, F. (2007), page 1ff

⁶⁹ See Itoh, H. (2004), page 18ff

⁷⁰ See Lammers, F (2010), page 169ff

⁷¹ See Lammers, F. (2007), page 1ff and Lammers, F. (2010), page 169ff

⁷² See Lammers, F. (2007), page 1ff and Lammers, F. (2010), page 169ff

Furthermore, according to the model of Itoh⁷³ if the project's output is small, the principal benefits from having a selfish agent as he is motivated to increase the possibility of raising the wage payment by choosing the high effort e_1 .

Offering low-powered incentives will therefore be enough for the selfish agent. If the agent compares his "net" income, thus including cost c , the principal can benefit from having an inequity averse agent as the incentive compatibility is easier to satisfy.

Actually in this model purely altruistic agents were not taken into account. So the fair or efficiency-seeking agent would have the parameters $e < 0$ and $\gamma = -1$. So the agent's utility would therefore be increasing in the principal's payoff and his income. If $\gamma = -1$, the agent is purely altruistic, where his utility is identical whether he is behind or ahead. If $\gamma < -1$ the agent stresses his income more and the principal's payoff less when he is behind than when he is ahead. So the principal benefits from a more altruistic agent, as the agent has an incentive to choose the high effort and so the principal can save the monetary incentive. If the agent is sufficiently altruistic, no monetary incentive is needed and only the participation constraint binds. As in both models the case of purely altruistic agents is excluded as this is rather the case of voluntary work.⁷⁴

⁷³ See Itoh, H., (2004), page 18ff

⁷⁴ See Itoh, H., (2004), page 18ff

6 Interaction with a third party

6.1 Fairness in delegated bargaining⁷⁵

6.1.1 Interaction with an external third party

6.1.1.1 Introduction

In fact this topic is not of great interest as Itoh⁷⁶ doesn't consider this setting in his model. But for better understanding and comparison with the next setting – internal third party – the scaling factor h has to be explained.

6.1.1.2 Scaling factor: Third party's payoff less sensitive to agent's effort

One can imagine that the external third party is not always a third party in usual case as a buyer but for example a supplier of the firm – so there is a kind of relation. Consider first the external third party's payoff function including the scaling factor:

$$\pi_B = h(1 - e)\theta_H \text{ with } h < 1$$

The effect of the agent's effort on his payoff is then reduced, as it could represent his anxiety to ensure a high quality of supplies for instance. Namely, the increase in quality costs the supplier less than it benefits the principal.

In fact the introduction of h is seen as a decline in the caring of the fair agent for the third party from μ_B to $h\mu_B$. The effort level and commission rate then are as follows:

$$e^{e,h}(F) = \frac{\beta + \mu_P(1 - \beta) - \mu_B h \theta_H}{c}$$
$$\beta^{e,h}(F) = \frac{1 - 2\mu_P + h\mu_B}{2(1 - \mu_P)}$$

⁷⁵ See Lammers, F. (2007), page 12ff

⁷⁶ See Itoh, H. (2004), page 18ff

Therefore the principal's condition to prefer a selfish agent changes to

$$1 - \mu_P > (1 - h\mu_B)^2$$

If $\mu_P = \mu_B$ one receives

$$\pi_P^{e,h}(S) > \pi_P^{e,h}(F) \leftrightarrow \mu_P < \frac{2h - 1}{h^2}$$

If $\mu_P \in (0; \frac{1}{2})$ the condition is fulfilled for $h > 2 - \sqrt{2}$ but it is not fulfilled for $h < \frac{1}{2}$.

For high values of h the agent's effect plays an important role for the third party's payoff. In comparison to a low h , where the effect of the agent's effort on the third party is reduced the fair agent becomes more attractive for the principal.

6.1.2 Interaction with an internal third party

6.1.2.1 The model

Again only the selfish respectively fair agent is able to exert effort. The main difference in this setting is the change in principal's preferences as he now employs both agents – the fair/selfish one and the selfish third party.

The timing is of course as before. This setting can have different internal interactions, meaning for example team production⁷⁷, helping on the job⁷⁸, and so forth.

Generally said the principal will always prefer a fair agent, if his effort increases his colleague's production, as the principal now has access to the whole production. From this positive effect for the principal and his colleagues, the fair agent feels additional utility.

Otherwise, the negative effect should be analyzed as the higher the agent's effort is, consequently the higher his own output q_A , but unfortunately this decreases to

⁷⁷ See Holmström, B. (1982), page 326

⁷⁸ See Drago, R., Garvey, G.T. (1998), page 1ff

some extent h the third party's (colleague's) result q_B . So total production then is given as

$$q = q_A + q_B$$

$$q = e\theta_H + h(1 - e)\theta_H, \text{ with } h \in [0,1]$$

In the extreme case, where $h = 0$, total production q equals the case of individual work. For $h = 1$ the principal won't encourage the agent in the case of internal interaction to exert effort, as aggregate surplus is maximized for $e = 0$ and chooses therefore $\beta = 0$.

So in this case it is more interesting to analyze $h \in (0,1)$.

In the case of a high h , after the surplus has been realized, there is an inefficient bargaining over the split between the agent and the third party. Therefore the agent's extent of effort is irrelevant for the principal.

In the case of low h , there are different situations considered, but according to Lammers⁷⁹, the agent's effort raises total production as the principal can imply cost savings for example.

Both agents, A and B are employed and receive a wage according to their individual production q_A and q_B which are in fact verifiable. The following assumption made is especially identified in larger firms, where it is too complex to introduce different payments based on the agent's effect on each employee. The third party has to bear the negative consequences of a decrease in production, although it was not due to his effort⁸⁰.

⁷⁹ See Lammers, F. (2007), page 20

⁸⁰ See Drago, R, Garvey, G.T. (1994), page 4

Again the agents receive a linear wage of

$$w_A = \alpha + \beta q_A \text{ and } w_B = \delta + \gamma q_B.$$

$$\begin{aligned} \pi_P &= q_A + q_B - w_A - w_B \rightarrow \\ e\theta_H + h(1-e)\theta_H - (\alpha + \beta e\theta_H) - (\delta + \gamma h(1-e)\theta_H) \end{aligned}$$

$$\begin{aligned} \pi_A &= w_A - \frac{c}{2}e^2 \rightarrow \alpha + \beta q_A - \frac{c}{2}e^2 \rightarrow \\ &\alpha + \beta e\theta_H - \frac{c}{2}e^2 \end{aligned}$$

and

$$\pi_B = w_B = \delta + \gamma q_B \rightarrow \delta + \gamma h(1-e)\theta_H$$

are the corresponding payoff functions for the principal, the fair respectively selfish agent and the third party. The internal third party receives the same commission rate ($\gamma = \beta$) and fixed payment ($\delta = \alpha$) as the agent, so as to rule out principal's preferences for a specific type of agent. Otherwise he can be interested to save wage payments as $h = 1$ his payoff function changes to

$$\begin{aligned} \pi_P &= q_A + q_B - w_A - w_B \rightarrow \\ e\theta_H + (1-e)\theta_H - (\alpha + \beta e\theta_H) - (\delta + \gamma(1-e)\theta_H) \rightarrow \\ &(\gamma - \beta)e\theta_H + (1 - \gamma)\theta_H - \alpha - \delta \end{aligned}$$

6.1.2.2 Choice of effort

For the selfish agent the introduction of a third party has no influence on his effort which depends on his commission rate

$$e^e(S) = e^i(S) = \beta \frac{\theta_H}{c}$$

so it is only of importance to consider the fair agent choices.

As he also considers the wage payment of his colleague, which is an add-on to his utility, he makes following effort decisions:

$$\begin{aligned}
U_F &= \alpha + \beta e^I(F) \theta_H - \frac{c}{2} e^I(F)^2 + \mu_P \pi_P + \mu_B \pi_P \\
&\rightarrow \alpha + \beta e^i(F) \theta_H - \frac{c}{2} e^i(F)^2 \\
&+ \mu_P \left((e^i(F) \theta_H + h(1 - e) \theta_H) - (\alpha + \beta e^i(F) \theta_H) - (\delta + \gamma h (1 - e^i(F)) \theta_H) \right) \\
&+ \mu_B (\delta + \gamma h (1 - e^i(F)) \theta_H)
\end{aligned}$$

Again after derivation one gets the effort of the fair agent.

$$e^I(F) = (\beta + \mu_P(1 - \beta - h(1 - \gamma)) - \mu_B \gamma h) \frac{\theta_H}{c}$$

And for equal commission rates $\gamma = \beta$ the effort changes to

$$e^i(F) = (\beta + \mu_P(1 - \beta)(1 - h) - \mu_B \beta h) \frac{\theta_H}{c}$$

Now the choice of effort also depends on the spillover factor h . So for a high h (inefficient haggling) the agent chooses a lower effort as a high effort would be disadvantageous to the principal and the internal third party ($\frac{\delta e^i(F)}{\delta h} < 0$).

If h is now introduced in the effort decision it is rewritten from

$$e^I(F) = (\beta + \mu_P(1 - \beta - h(1 - \gamma)) - \mu_B \gamma h) \frac{\theta_H}{c}$$

to

$$e^i(F) = (\beta + \mu_P(1 - \beta) - \mu_B h + h(1 - \beta)(\mu_B - \mu_P)) \frac{\theta_H}{c}$$

If the agent is equally interested in the principal's and third party's payoff

$$\mu_P = \mu_B$$

the effort changes to

$$e^i(F) = (\beta(1 - \mu_P) + \mu_P(1 - h)) \frac{\theta_H}{c}$$

If the agent cares more for the principal

$$\mu_P > \mu_B$$

the fair agent exerts less effort in the internal setting and he invests more effort in the case of

$$\mu_P < \mu_B$$

In other words, if he cares more for the third party, so his colleague, the fair agent exerts more effort. The reason is because this doesn't endamage the third party as much, since the third party only receives β of his production. So effort now, to some extent, also harms the principal.

6.1.2.3 Optimal employment contracts

The principal maximizes his payoff by taking the wage of the third party as given. For the selfish agent the following payoff and utility function are of interest:

$$\begin{aligned} \max \pi_P &= q_A + q_B - w_A - w_B \\ &\rightarrow e^i(S)\theta_H + h(1 - e^i(S))\theta_H - (\alpha + \beta e^i(S)\theta_H) \\ &\quad - (\delta + \gamma h(1 - e^i(S))\theta_H) \\ &\rightarrow (1 - \beta)e^i(S)\theta_H + (1 - \gamma)h(1 - e^i(S))\theta_H - \alpha - \delta \end{aligned}$$

$$\text{s.t. } U_S = \pi_A = \alpha + \beta e^i(S) \theta_H - \frac{c}{2} e^i(S)^2 \rightarrow \alpha + \beta \left(\beta \frac{\theta_H}{c} \right) \theta_H - \frac{c}{2} \left(\beta \frac{\theta_H}{c} \right)^2 \rightarrow$$

$$\alpha + \frac{(\beta \theta_H)^2}{2c} \geq 0 \text{ and } \alpha \geq 0$$

As the limited liability constraint again has to and the participation constraint cannot be binding and after using the Lagrange function, the selfish agent will not receive a fixed payment

$$\alpha^i(S) = 0$$

and a commission rate

$$\beta^i(S) = \frac{1 - h(1 - \gamma)}{2}$$

The agent and the internal third party receive the same commission rates

$$\beta = \gamma \text{ and fixed payments } \alpha = \delta$$

so to insure that the fairness preferences of the agents are the decision factor for principal's hiring decisions and not perhaps the wage differences.

$$\beta^i(S) = \frac{1 - h(1 - \gamma)}{2} \rightarrow \frac{1 - h(1 - \beta)}{2} \rightarrow \frac{1 - h}{2 - h}$$

$$\begin{aligned} \pi_P^i(S) &= (1 - \beta) e^i(S) \theta_H + (1 - \gamma) h \left(1 - e^i(S) \right) \theta_H - \alpha - \delta \\ &\rightarrow (1 - \beta) \left(\beta \frac{\theta_H}{c} \right) \theta_H + (1 - \gamma) h \left(1 - \left(\beta \frac{\theta_H}{c} \right) \right) \theta_H - \alpha - \delta \\ &\rightarrow \left(1 - \left(\frac{1 - h}{2 - h} \right) \right) \left(\left(\frac{1 - h}{2 - h} \right) \frac{\theta_H}{c} \right) \theta_H + (1 - \gamma) h \left(1 - \left(\left(\frac{1 - h}{2 - h} \right) \frac{\theta_H}{c} \right) \right) \theta_H \\ &\quad - \alpha - \delta \rightarrow \end{aligned}$$

$$\pi_P^i(S) = \frac{(1 - h)^2 \theta_H^2}{(2 - h)^2 c} + \frac{h}{2 - h} \theta_H$$

The fair agent will receive a contract solving following maximization problem:

$$\begin{aligned}
\max \pi_P &= q_A + q_B - w_A - w_B \\
&\rightarrow e^i(F)\theta_H + h(1 - e^i(F))\theta_H - (\alpha + \beta e^i(F)\theta_H) \\
&\quad - (\delta + \gamma h(1 - e^i(F))\theta_H) \\
&\rightarrow (1 - \beta)e^i(F)\theta_H + (1 - \gamma)h(1 - e^i(F))\theta_H - \alpha - \delta
\end{aligned}$$

$$\begin{aligned}
\text{s.t. } U_F &= \alpha + \beta e^i(F)\theta_H - \frac{c}{2}e^i(F)^2 + \mu_P\pi_P + \mu_B\pi_B \geq 0 \rightarrow \alpha + \beta e^i(F)\theta_H - \frac{c}{2}e^i(F)^2 + \\
&\quad \mu_P \left((1 - \beta)e^i(F)\theta_H + (1 - \gamma)h(1 - e^i(F))\theta_H - \alpha - \delta \right) + \mu_B \left(\delta + \gamma h(1 - \right. \\
&\quad \left. e^i(F))\theta_H \right) \geq 0 \text{ and } \alpha \geq 0
\end{aligned}$$

To keep it well-structured, the above formula is left short. Again, after using the Lagrange-function, in fact the agent will receive no fixed payments

$$\alpha^i(F) = 0$$

and a commission rate

$$\beta^i(F) = \frac{1 - 2\mu_P - h(1 - \gamma) + 2h\mu_P(1 - \gamma) + h\gamma\mu_B}{2(1 - \mu_P)}$$

Again both, the agent and the third party, receive the same commission rate

$$\beta^i(F) = \gamma^i(F)$$

The principal's payoff function with a fair agent is therefore

$$\begin{aligned}
\pi_P^i(F) &= (1 - \beta)e^i(F)\theta_H + (1 - \gamma)h(1 - e^i(F))\theta_H - \alpha - \delta \rightarrow \\
\pi_P^i(F) &= \frac{(1 - \mu_P)(1 - h\mu_B)^2(1 - h)^2}{(2(1 - \mu_P) - h(1 - 2\mu_P + \mu_B))^2} \frac{\theta_H^2}{c} + \frac{h(1 - h\mu_B)}{2(1 - \mu_P) - h(1 - 2\mu_P + \mu_B)} \theta_H
\end{aligned}$$

As the agent of course requires some monetary compensation, the commission is positive for $\mu_P < 1$. $\frac{\delta \pi_P^i(F)}{\delta \mu_P} > 0$ and $\frac{\delta \pi_P^i(F)}{\delta \mu_B} < 0$, so the more the agent cares for the principal and the less he cares for the third party, the higher the principal's payoff. Furthermore $\frac{\delta \beta^i(F)}{\delta h} < 0$ and $\frac{\delta \beta^{e,h}(F)}{\delta h} > 0$, so in the internal setting, a higher h lowers the effort of the agent and therefore the principal pays a lower commission rate. So to induce the agent to spend more effort, the principal has to pay a higher commission rate the more the actions negatively affect the third party.

6.1.2.4 Choice of agent

The principal prefers to hire a fair agent of course if

$$\begin{aligned} \pi_P^i(F) > \pi_P^i(S) &\rightarrow \\ \frac{(1 - \mu_P)(1 - h\mu_B)^2(1 - h)^2}{(2(1 - \mu_P) - h(1 - 2\mu_P + \mu_B))^2} \frac{\theta_H}{c} + \frac{h(1 - h\mu_B)}{2(1 - \mu_P) - h(1 - 2\mu_P + \mu_B)} \\ &> \frac{(1 - h)^2}{(2 - h)^2} \frac{\theta_H}{c} + \frac{h}{2 - h} \end{aligned}$$

This condition holds if the concern for the principal and the third party is equal ($\mu_P = \mu_B$). Clearly, if the fair agent only cares of the third party (μ_B) then the principal will prefer the selfish agent.

The principal often hires a fair agent in internal settings if

$$\pi_P^i(F) > \pi_P^i(S),$$

because the principal participates to some extent $(h(1 - e)(1 - \beta)\theta_H)$ from the production of the third party, so the concern somehow benefits the principal.

Two cases are now considered: if the agent has equal concern for the principal and the third party and the case if the agent cares more for the third party.

6.1.2.4.1 Same concern for both

If he cares equally for the principal and the third party ($\mu_P = \mu_B$) one gets from $\pi_P^i(F) > \pi_P^i(S)$

$$\mu_P > \frac{2h-1}{h^2} - \frac{(2-h)c}{h(1-h)\theta_H}$$

This condition holds all the time. In this case the principal will always hire a fair agent because for low h the principal wants a high effort level. Therefore the fair agent will take the principal's payoff strongly into account as this doesn't endamage the third party's payoff too much. So the fair agent is cheaper. For a high h the concerns for the third party also benefit the principal's payoff and so the fair agent is again cheaper to employ.

6.1.2.4.2 No concern for principal but concern for third party

If the agent has a higher concern for the third party and no concern for the principal (e.g. no contact to the management, close contact to colleagues) this induces the following, which can also be seen in Figure 6.

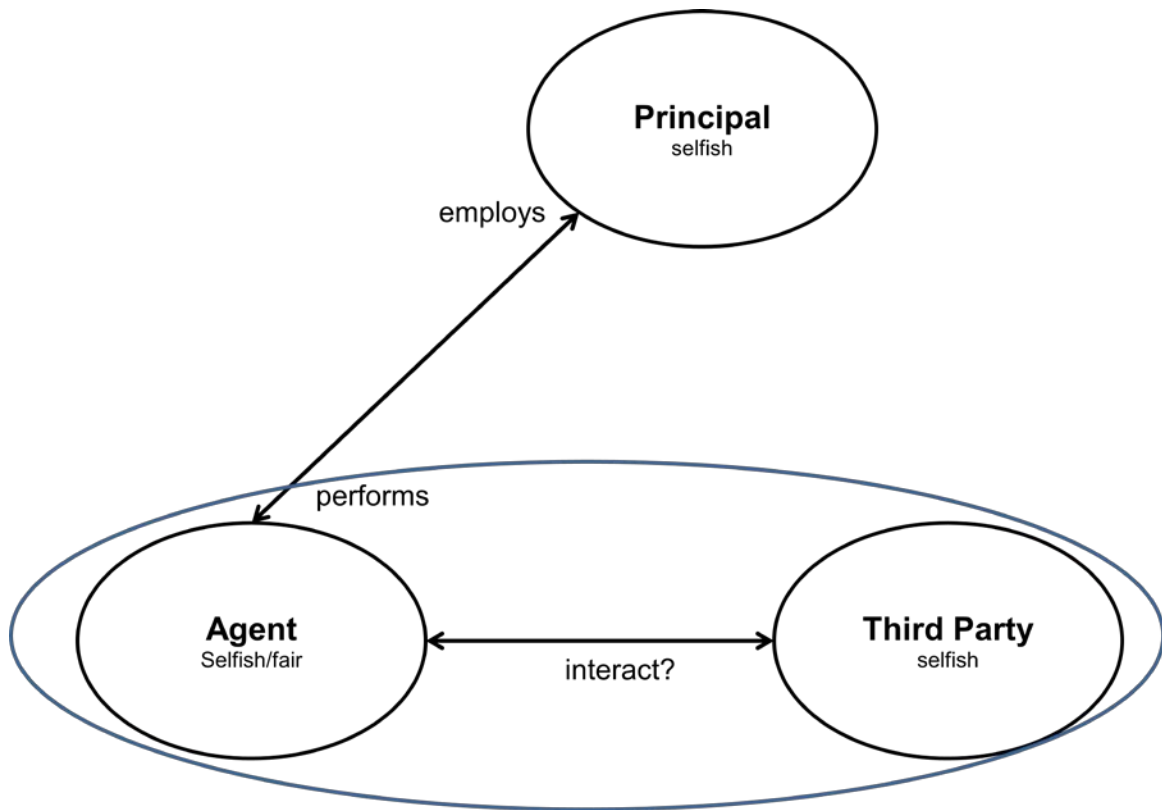


Figure 6: Only concern for third party⁸¹

As the selfish agent cares only for his own monetary payoff, his effort remains the same. As a consequence the optimal employment contract is the same.

In opposite the fair agent considers also the wage payment of the third party and his choice of his effort also depends on the spillover factor h . So for a high h the agent chooses a lower effort as a high effort would be disadvantageous to the internal third party.

The principal then prefers the fair agent if

$$\pi_p^i(F) > \pi_p^i(S) \Leftrightarrow$$

⁸¹ Own illustration

$$-h > \frac{(1-h)^2((2-h\mu_B)(1-h) + 2(1-h\mu_B))}{(2-h)(2-h(1+\mu_B))} \frac{\theta_H}{c}$$

As the left hand side is negative and the right hand side is positive, this condition can never be fulfilled. The fair agent will therefore exert no effort, so the selfish agent will be preferred. The principal will again have an incentive to establish and strengthen the fair agent's concern for himself.

6.1.2.4.3 Influence of the spillover factor

Given $\frac{\delta\pi_P^i(F)}{\delta\mu_P} > 0$ and $\frac{\delta\pi_P^i(F)}{\delta\mu_B} < 0$ and an assumed value $\frac{4}{3} = \frac{\theta_H}{c}$.

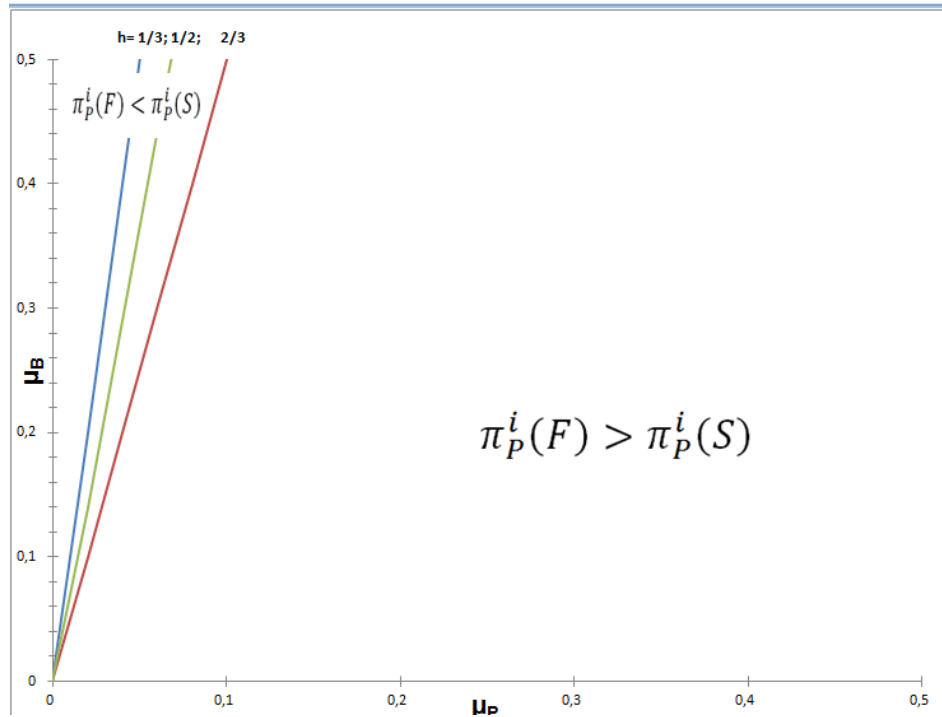


Figure 7: Choice of agent in the internal setting⁸²

⁸² Lammers, F. (2007), page 26

As can be seen from the Figure 7 above, the principal will more often prefer a fair agent, especially in the case of a low spillover factor h . With a low h the third party suffers only little from the agent's action as with a low h the concern is signified as being less relevant.

6.2 Moral hazard and other-regarding agents⁸³

6.2.1 The model

In this setting each agent cares about the material payoff of the other agents, so the third party, but not about the principal's payoff as can also be seen in Figure 6. The principal employs two agents, who work separately and are denoted by $n = 1, 2$.

The timing is again the same and the payment is the following:

$$w^n = (w_{jk}^n) \quad \text{for } j, k = s, f$$

w_{jk}^n denotes the wage payment to the agent n when his outcome is j and the other agent's outcome is k . Again the payment scheme must satisfy the limited liability constraints

$$w_{jk}^n \geq 0, j, k = s, f$$

Agent's n utility function is the following

$$u^n = w_{jk}^n - c_i - \mu_{B_n} v_n(\max\{w_{kj}^m - w_{jk}^n, 0\}) - \mu_{B_n} \gamma_n v_n(\max\{w_{jk}^n - w_{kj}^m, 0\})$$

for $i = 0, 1$, $j, k = s, f$, $n, m = 1, 2$ and $m \neq n$. Again $\mu_{P_n} \geq 0$, $v_n(0) = 0$, $v'_n(z) > 0$ for $z > 0$. In this case the experimental evidence that $|\gamma_n| \leq 1$ holds, so the

⁸³ See Itoh, H. (2004), page 32ff

agents suffer more from being behind than being ahead. Being behind or ahead by a certain amount changes the utility by an equal amount.

The agents compare their wages for each pair of project outcomes, because each agent observes or estimates what the other agent is paid, while the other agent's choice of action is not a concern.

It is now assumed (**Assumption 2**) that

- a) the agents are symmetric: $\mu_B = \mu_{B_1} = \mu_{B_2}$, $\gamma = \gamma_1 = \gamma_2$ and $v(\cdot) = v_1(\cdot) = v_2(\cdot)$
- b) the principal offers a symmetric contract that satisfies $w^1 = w^2$
- c) $v(z) = z$ for all $z \geq 0$
- d) $\mu_B \gamma \leq 1$

These assumptions conform with the assumptions of the previous model. If d) doesn't hold, the inequity-averse agent, when he is ahead, wants to reduce the inequality by giving up his income (which is seen as voluntary work). But this seems implausible. The agent's utility function can thus be rewritten as

$$u^n = \begin{cases} w_{jk}^n - c_i - \mu_B \gamma (w_{jk}^n - w_{kj}^m) & \text{if } w_{jk}^n \geq w_{kj}^m \\ w_{jk}^n - c_i - \mu_B (w_{kj}^m - w_{jk}^n) & \text{if } w_{jk}^n \leq w_{kj}^m \end{cases}$$

As before, the principal's benefit q is large enough to implement (e_1, e_1) . So as to reduce the expected payments to the agents, he will offer a symmetric contract $(w^1, w^2) = (w, w)$. As the principal pays the least possible amount to each agent when his project fails: $w_{fs} = w_{ff} = 0$ - so if at least one project fails, no one receives any payment. The incentive compatibility constraint and the participation constraint are the following:

$$p_1 w_{ss} + (1 - p_1) w_{sf} + [p_1 - (1 - p_1) \gamma] \mu_B w_{sf} \geq \frac{c}{\Delta_p} \text{ (IC)}$$

$$p_1 w_{ss} + (1 - p_1) w_{sf} - (1 - p_1) \mu_B (1 + \gamma) w_{sf} \geq \frac{c}{p_1} \text{ (PC)}$$

One can see, if the project of one agent fails and the project of the other agent succeeds, the agent whose project fails is behind and in opposite the agent whose project succeeds is ahead. So the agent who is behind suffers $\mu_B w_{sf}$ from inequity aversion and the agent who is ahead, his equity concern is represented by $\mu_B \gamma w_{sf}$.

If $w_{ss} = w_{sf}$ w is called an independent contract as the contract depends only on the outcome of the agent's individual project. Compared to $w_{ss} > w_{sf}$ it is a team contract and $w_{ss} < w_{sf}$ is called a relative performance contract. Of course the principal chooses the contract that minimizes his expected payments as a function of the participation constraint and incentive compatibility constraint

$$2(p_1 w_{ss} + (1 - p_1) w_{sf}) \rightarrow 2(p_1 w_{ss} + w_{sf} - p_1 w_{sf}) \rightarrow$$

$$2p_1 (w_{sf} + p_1 (w_{ss} - w_{sf}))$$

6.2.2 The self-interest case

The benchmark case of purely self-interested agents, so $\mu_B = 0$, is analyzed. In this case it is clear that only the incentive compatibility constraint binds and the contract is optimal that solves

$$p_1 w_{ss} + (1 - p_1) w_{sf} = \frac{c}{\Delta_p}$$

The principal's expected payment to the agents is therefore

$$2p_1 (w_{sf} + p_1 (w_{ss} - w_{sf})) \rightarrow 2p_1 (w_{sf} (1 - p_1) + p_1 w_{ss}) \rightarrow 2p_1 \left(\frac{c}{\Delta_p} \right) \rightarrow$$

$$\frac{2p_1 c}{\Delta_p}$$

It has to be noted here that the independent contract is optimal, where

$$w_{ss} = w_{sf} = \frac{c}{\Delta_p}$$

6.2.3 Other-regarding agents

It is first supposed that the rearranged incentive compatibility constraint is binding

$$p_1 w_{ss} + (1 - p_1) w_{sf} \geq \frac{c}{\Delta_p} + [(1 - p_1)\gamma - p_1] \mu_B w_{sf}$$

The larger the left-hand side the higher is the principal's expected payment. So if $(1 - p_1)\gamma > p_1$ from the right-hand side, the principal will choose a team contract $w_{ss} > w_{sf} = 0$ (Case 1) compared to the case if $(1 - p_1)\gamma < p_1$, the principal will choose a relative performance contract $w_{ss} = 0 < w_{sf}$ (Case 2) so as to satisfy this inequality. If $(1 - p_1)\gamma = p_1$ any contract which satisfies the incentive compatibility constraint is optimal.

The optimal $\hat{w}_{ss}$ and $\hat{w}_{sf}$ are defined by derivation of the incentive compatibility constraint and are

$$\hat{w}_{ss} = \frac{c}{\Delta_p} \frac{1}{p_1}$$

$$\hat{w}_{sf} = \frac{c}{\Delta_p} \frac{1}{(1 - p_1) + \mu_B [p_1 - (1 - p_1)\gamma]}$$

For simplicity the two cases are analyzed separately:

(i) Case 1: $(1 - p_1)\gamma > p_1$

The best contract in this case is $(\hat{w}_{ss}, 0)$. This contract also satisfies the participation constraint by $p_1 > \Delta_p$ and is therefore the optimal contract. The expected payment equals the case with only purely self-interested agents and is therefore

$$\frac{2p_1c}{\Delta_p}$$

(ii) Case 2: $(1 - p_1)\gamma < p_1$

The incentive compatibility constraint binds at $(0, \hat{w}_{sf})$ and one needs to derive the participation constraint by substituting $(0, \hat{w}_{sf})$ into the participation constraint (PC):

$$p_1 w_{ss} + (1 - p_1) w_{sf} - (1 - p_1) \mu_B (1 + \gamma) w_{sf} \geq \frac{c}{p_1} \rightarrow$$

$$\frac{c}{\Delta_p} \frac{(1 - p_1)(1 - \mu_B(1 + \gamma))}{(1 - p_1) + \mu_B[p_1 - (1 - p_1)\gamma]} \geq \frac{c}{p_1}$$

The following likelihood ratios are needed for a straightforward calculation:

$$l_s = \frac{p_1}{p_0}$$

$$l_f = \frac{(1 - p_1)}{(1 - p_0)}$$

which is needed to satisfy the participation constraint:

$$\frac{l_s}{l_f} \leq \frac{1 - \mu_B \gamma}{\mu_B}$$

So when the conditions of Case 2 and the likelihood ratios hold then $(0, \hat{w}_{sf})$ is the optimal contract which is called Case 2a.

The expected payment then is

$$2p_1(1 - p_1)\hat{w}_{sf} = \frac{2c}{\Delta_p} \frac{p_1(1 - p_1)}{(1 - p_1) + \mu_B[p_1 - (1 - p_1)\gamma]}$$

One can see that this is equal to the Case 1 $\frac{2p_1c}{\Delta_p}$ if $\mu_B = 0$. Furthermore this equation is decreasing in μ_B and increasing in γ .

Nevertheless if $(0, \hat{w}_{sf})$ cannot satisfy the straightforward calculation, the participation constraint must also bind:

$$p_1 w_{ss} + (1 - p_1) w_{sf} = \frac{c}{p_1} + (1 - p_1) \mu_B (1 + \gamma) w_{sf}$$

The larger the left-hand side, the larger the principal's payment is. The principal will therefore want to choose w_{sf} as small as possible. Both, the incentive compatibility and the participation constraints bind as $w_{sf} = 0$ doesn't satisfy the incentive compatibility constraint. Solving the equations yield

$$\bar{w}_{ss} = \frac{c}{\Delta_p} \frac{1 - p_1}{l_s p_1} \left(\frac{l_s}{l_f} - \frac{1 - \mu_B \gamma}{\mu_B} \right)$$

$$\bar{w}_{sf} = \frac{c}{\Delta_p} \frac{1}{l_s \mu_B}$$

If the straightforward calculation doesn't hold, $\bar{w}_{ss} > 0$. The expected payment is then

$$2p_1 [p_1 \bar{w}_{ss} + (1 - p_1) \bar{w}_{sf}] = \frac{2c p_1 (1 - p_1)}{\Delta_p l_s} \left(\frac{l_s}{l_f} + \gamma \right)$$

One can see that the payment is independent of μ_B and increasing in γ . This case is called Case 2b. This case occurs if $\mu_B^{-1} - \left(\frac{l_s}{l_f} \right) < \gamma < \frac{p_1}{(1 - p_1)}$, a range which exists only if $\mu_B > \frac{(1 - p_1)p_0}{p_1}$.

There occur two incentive effects via the incentive compatibility constraint if there are other-regarding preferences.

- First, if the agent is behind because his project fails while the other agents succeed, he feels disutility $\mu_B w_{sf}$ which brings about a positive incentive effect (works more).
- Secondly, if the agent is ahead because his project succeeds while the other's fail, the inequity-averse agent suffers if he is ahead. Therefore he is not willing to increase the probability of success, which is negative for inequity-averse agents.

- o In Case 1 this negative effect dominates because each agent is averse enough being ahead of each project is uncertain and uncontrollable so that the incentive effect of being ahead seems more important than that of being behind.

Since the effects of inequity aversion should be minimized in this case, the extreme team contract is used. So the agents receive the same payoff, so for the principal it is equal if they care about each other's wellbeing, so the principal neither benefits nor suffers from the other-regarding preferences.

- o If it is now supposed that the project is controllable or the extent to which the agent is averse being ahead is small, the first positive incentive effect dominates the second one. This is in our setting Case 2a or Case 2b.

The principal will therefore use a relative performance contract to create the possibility of inequity. So if the agents are status seeking the first and the second effect are positive, as they prefer being ahead. A tournament-like extreme relative performance contract seems to be optimal as long as the participation contract doesn't bind (Case 2a). So the introduction of "competition" doesn't benefit the principal if the agents are self-interested.

Nevertheless, the principal has to compensate the disutility from other-regarding preferences. In Case 2b the participation constraint binds and the principal offers an extreme relative performance contract and chooses contract in which the agents are paid a positive amount whether the project succeeds or fails. If now the optimal contract in case of purely self-interested agents is

compared with the contract mentioned above, one sees that the independent contract in the case of self-interested agents is optimal. If there are other-regarding preferences this contract seems no longer to be optimal, despite technological and stochastic independence.

6.2.4 “Net” Material Payoff

In this section the agents compare their net material payoff, as for example agents working closely are able to observe their actions or can expect the other's actions correctly. One can imagine that the actions will be compared. The agent's utility function is now the following:

$$u^n = \begin{cases} w_{jk}^n - c_i - \mu_B \gamma (w_{jk}^n - c_i - w_{kj}^m + c_h) & \text{if } w_{jk}^n - c_i \geq w_{kj}^m - c_h \\ w_{jk}^n - c_i - \mu_B (w_{kj}^m - c_h - w_{jk}^n + c_i) & \text{if } w_{jk}^n - c_i \leq w_{kj}^m - c_h \end{cases}$$

Again, $j, k = s, f$, $h, i = 0, 1$ and i is the index for agent n and h represents the index for agent m . If both agents choose e_1 , their utility doesn't change from the previous model and the participation constraint is thus equal. In contrast, if one agent chooses e_o and the other e_1 , the comparison is now between w_{jk}^n and $w_{kj}^m - c_h$. The incentive compatibility constraint is therefore

$$p_1 w_{ss} + (1 - p_1) w_{sf} + [p_1 - (1 - p_1) \gamma] \mu_B w_{sf} \geq (1 - \mu_B \gamma) \frac{c}{\Delta_p} + (1 - p_0) p_1 \mu_B (1 + \gamma) \frac{w_{sf}}{\Delta_p} \text{ if } w_{sf} \leq c$$

$$p_1 w_{ss} + (1 - p_1) w_{sf} + [p_1 - (1 - p_1) \gamma] \mu_B w_{sf} \geq [(1 - \mu_B \gamma) + (1 - p_0) p_1 \mu_B (1 + \gamma)] \frac{c}{\Delta_p} \text{ if } w_{sf} \geq c$$

If first an extreme team contract is considered, the incentive compatibility constraint is

$$p_1 w_{ss} \geq (1 - \mu_B \gamma) \frac{c}{\Delta_p}$$

The agent is therefore always ahead by c if he chooses e_0 and the other chooses e_1 . If he is status-seeking the agent will enjoy this divergence and the principal has therefore to offer stronger incentives to induce him to exert more effort. Otherwise, if the agent is inequity averse, he will choose e_1 so as to avoid being ahead and the principal can thus save on the payments.

The principal is therefore better off because these agents feel bad if they shirk. This non-pecuniary incentive substitutes the monetary incentive. The principal's payment now decrease in μ_B under the extreme team contract, thus the principal wants the agents to be more inequity averse if the extreme team contract is used.

This result differs from the previous one, where the agent is indifferent when implementing the extreme team contract. Moreover, if the agents are sufficiently inequity averse, the participation constraint becomes binding and the principal needs not to leave rents to them: the participation constraint is therefore

$$p_1 w_{ss} \geq \frac{c}{\Delta_p}$$

and equals the previous incentive compatibility constraint only if $\mu_B \gamma \geq \frac{1}{l_s}$.

This new incentive benefits also exist under relative performance contract, but the benefit is not as large as under team contracts. As $\hat{w}_{sf} > c$ the incentive compatibility constraint applies, where the right-hand side is larger than that in the case of extreme team contracts if $\gamma > -1$. If one project fails and the other succeeds, the agent whose project fails is behind. So if he is shirking, he compares his payoff zero with the other agent's payoff $\hat{w}_{sf} > 0$. But the extent to which the agent is behind is not as important as in the case of unobservable actions, so the new specification provides a negative incentive. So team contracts may be more attractive, and in fact more optimal, for the principal if agents compare their actions and their incomes.

Imagine if (i) Case 1 $(1 - p_1)\gamma > p_1$ holds, the extreme team contract is optimal and if (ii) $\mu_B\gamma \geq \frac{1}{l_s}$ holds the extreme team contract where the participation constraint binds is optimal.

- (i) It is supposed that a contract with $w_{sf} > 0$ is optimal. w'_{ss} is defined by

$$p_1 w'_{ss} = p_1 w_{ss} + (1 - p_1)w_{sf} - \varepsilon$$

and assume an extreme team contract where $\varepsilon > 0$. One sees that the incentive compatibility constraint and the participation constraint for ε sufficiently small but positive are satisfied. The principal's payments are then

$$2p_1^2 w'_{ss} < 2p_1 [p_1 w_{ss} + (1 - p_1)w_{sf}]$$

which differs from the optimality of (w_{ss}, w_{sf}) .

- (ii) The principal's expected payment is $2c$. If now an arbitrary contract is assumed the principals' payments are

$$2p_1 [p_1 w_{ss} + (1 - p_1)w_{sf}] \geq 2p_1 \left(\frac{c}{p_1} + (1 - p_1)\mu_B(1 + \gamma)w_{sf} \right) \geq 2c$$

The first inequality follows from the participation constraint and the second inequality holds because it is assumed that $\gamma \geq -1$. So the extreme team contract is optimal.

Concluding it can be said that

- (i) an extreme team contract is optimal.
- (ii) In fact there exists a range of parameters in which the extreme team contract is optimal if the agents compare their "net" material payoff, although it is not in the original one. The conditions are $\mu_B\gamma \geq \frac{1}{l_s}$ and $(1 - p_1)\gamma < p_1$.

6.3 Comparison of “Interaction”

“Fairness in Delegated Bargaining”	“Moral hazard and other-regarding preferences”
Fair if at least same concern and selfish if no concern for principal	Rather selfish agent, but fair agent if cost of effort is included

Table 3: Comparison of "Interaction"

Table 3 compares the cases of multiple agents. The third party is for example a supplier and not a buyer, the agent fears to decrease the quality. An increase in quality is less costly for the supplier than it benefits the principal. As the agent's caring for the third party declines and in case of a low spillover, the fair agent becomes more attractive to the principal.⁸⁴

In case of internal interaction, the principal prefers a fair agent as long as the effort increases the colleague's (=third party) production as the principal now has access to the whole output. This positive effect for the principal and the third party increases the fair agent's utility as he feels additional utility due to the convexity of the function. This can be seen in the utility function where the fairness concern acts as an add-on.

But the higher the fair agent's output, the lower the third party's output as it is lowered to some extent h . As the effort choice now also depends on the spillover factor h , and h is high, the agent chooses a lower effort as this would harm the

⁸⁴ Lammers, F. (2010), page 169ff and Lammers, F. (2007), page 1ff

third party (decreases his output) and in turn the principal (payoff is decreased as total output is decreased).

If the agent cares to the same extent for the principal and the third party and in case of a low h the fair agent becomes more attractive to the principal. Therefore the choice of agent depends now also on the spillover factor. If the agent cares more for the principal, the fair agent exerts less effort (so not to decrease the third party's output) et vice versa. More effort doesn't harm the third party as much. So more effort now, to some extent, also harms the principal. Both agents, the selfish respectively fair agent and the third party, receive the same wage, so the same fixed payments and commission rate so as to rule out that the principal chooses an agent because of wage differences.

The more the agent cares for the principal and the less he cares for the third party, the higher is the principal's payoff as the agent exerts more effort. So again a high h decreases the agent's effort and the principal has to pay a lower commission rate. So the principal has to pay a higher commission rate to induce the agent to invest more effort the more the actions negatively affect the third party.

So in case of equal concern for the principal and the third party, the principal will choose the fair agent if his payoff is higher than by choosing a selfish agent.

If the agent cares only for the third party, the principal will prefer the selfish agent.

So in case of internal interaction, the principal often prefers the fair agent as the principal has access to the whole production, so the concern for the third party somehow benefits the principal. The principal wants a low h so as to induce the fair agent to exert high effort as it doesn't harm the third party too much. So the fair agent is cheaper to employ. Also a high h benefits the principal and so the fair agent is cheaper to employ.

If the agent only has concern for the third party, so his colleague, the agent won't exert any effort and thus the selfish agent will be preferred.⁸⁵ As can be seen in both models the principal decides whether to choose a fair or selfish agent depending on his payoff function and in turn on the concern factor μ_p which is included in the convex utility function.

In the model of Itoh⁸⁶ with multiple agents, as the agents compare their incomes, the principal wants the agents to be less inequity averse, so to have γ as small as possible, but μ_B - so the concern for him - as large as possible. Of course the principal wants a high concern for him. The agents compare their wages for each pair of outcomes. The principal offers a selfish agent an independent contract, as his contract depends only on his outcome. In the case of inequity averse agents, the offered contract depends on the fairness preference γ . So if the probability that the project fails and the social preference is higher than the probability that the project succeeds, the principal offers a team contract (Case 1).

In opposite, if the success of the project is higher than the probability that the project fails with fairness concern than the principal chooses a relative performance contract (Case 2). The principal's payment, thus agent's wage is increasing in γ . If only inequity-averse agents are feasible, then the optimal preference is $\gamma = 0$ and $\mu_B > \frac{l_f}{l_s}$, that is the agent doesn't feel guilty but envious (incentive to choose e_1). If also status-seeking agents are feasible, then the most competitive and sufficiently other-regarding agent is the best as in the Case 2b. There neither the extreme team nor the extreme relative performance contract was optimal. A reason why such competitive preference pattern are desirable is a lack of productive interaction amount the agents. So the principal neither benefits nor suffers from the agent's other-regarding preferences.

Furthermore the principal can change the preferences of the agents such that they compare also their actions and not only their payments. In this case the

⁸⁵ See Lammers, F. (2007), page 1ff

⁸⁶ See Itoh, H. (2004), page 32ff

principal will implement the extreme team contract with inequity-averse preferences as can be seen above. So, even if the agents work independently on their projects, an implementation of a firm, team-based pay scheme could benefit the principal if he can change the agent's preferences.⁸⁷

It can be seen that the main decision factor is the utility function, which is in case of Itoh benchmark-driven. So trespass a threshold decreases the utility and this in turn influences the decision for the fair or selfish agent. An interesting extension of Itoh's model would be to include fairness preferences for the principal as well in the utility function. Or how would the results change if the third party can also be either selfish or fair?

⁸⁷ Itoh, H. (2004), page 18ff

7 Conclusion

The standard economic theory can be disproved, as there are beside selfish also fair people caring for the wellbeing of others. As can be seen economists fail if they only believe on the self-interest hypothesis and rule out any social preferences.⁸⁸ Which agent – selfish or fair – is preferred depends on the individual assumptions made and was shown by analyzing two models:

- Lammers, F. (2010): “Fairness in Delegated Bargaining”⁸⁹
- Itoh, H. (2004): “Moral hazard and other-regarding preferences”⁹⁰

Generally the two models differ different cases:

- The agent doesn't interact with another party, so he works individually
- The agent interacts with a third party

Both models are described and analyzed as they have a lot of similarities – equal basic assumptions. It can be seen that in both models the principal decides whether to choose a fair or selfish agent depending on his payoff function and in turn on the concern factor μ_p . But they differ in a main point: The fairness consideration, as Itoh takes inequity averse agents into account and the utility function. Lammers uses a convex function, while Itoh uses a utility function where the optimum is at the break-even-point.

In the model of Lammers⁹¹ the convex (!) utility function depends on the wage, which includes a fixed and variable component. As a result the fair agent spends more effort as he cares for fairness concerns because this acts as an add-on to

⁸⁸ See Fehr, E., Fischbacher, U. (2002), page 1

⁸⁹ See Lammers, F. (2010), page 169ff

⁹⁰ See Itoh, H. (2004), page 18ff

⁹¹ See Lammers, F. (2007), page 1ff and Lammers, F. (2010), page 169ff

his utility. This means his effort raises the more he cares for the principal. So the principal has to pay the fair agent less which in turn increases his own payoffs.⁹²

In the model of Itoh⁹³ the agent compares his payoff with the principal's income and this turns out whether the agent is selfish ("ahead") or fair ("behind"). The principal has to pay the self-interested agent rents and leave him therefore high-powered incentives. In the single-agent case the principal generally does not benefit from the agent's other-regarding – social - preferences. This result is different from the model of Lammers⁹⁴. It is more costly for the principal to induce the other-regarding agent to exert more effort. This is because of the utility function in case of inequity aversion, where his utility decreases after a certain break-even-point.⁹⁵

In case of internal interaction, the principal prefers a fair agent as long as the effort increases the third party's (for example a colleague's) production as the principal now has access to the whole output. This positive effect for the principal and the third party increases the fair agent's utility as he feels additional utility due to convexity.⁹⁶

In the model of Itoh with multiple agents, as the agents compare their incomes, the principal wants the agents to be less inequity averse, but the concern for him as large as possible. The agents compare their wages for each pair of outcomes.

The principal offers a selfish agent an independent contract, as his contract depends only on his outcome. In the case of inequity averse agents, the offered contract depends on the fairness preferences. So if the probability that the project fails and the social preference is higher than the probability that the project succeeds, the principal offers a team contract. In opposite, if the success of the

⁹² See Lammers, F. (2010), page 169ff and Lammers, F. (2007), page 1ff

⁹³ See Itoh, H. (2004), page 18ff

⁹⁴ See Lammers, F. (2007), page 1ff and Lammers, F. (2010), page 169ff

⁹⁵ See Itoh, H. (2004), page 18ff

⁹⁶ See Lammers, F. (2007), page 18ff

project is higher than the probability that the project fails with fairness concern than the principal chooses a relative performance contract. The principal's payment is increasing in the inequity aversion.⁹⁷

Of course models, especially mathematical ones, are a simplified description of a system to describe it in an easy manner. As can be seen, results depend on the individual assumptions made. Although a lot of models considering social preferences emerged in recent years, there is no clear outcome observable. There are different models with different assumptions made. The two models in this thesis equal in many assumption. But the utility function is the factor which influences the results in a deeper way as they differ basically. An interesting point would be to try to adjust the utility functions to have a basic starting point. How would the result change if it is a repeated game? Would the principal change his opinion and choose the other agent? So including a learning curve would be another adjustment.

⁹⁷ Itoh, H. (2004), page 18ff

8 Summary

8.1 English

As can be seen economists fail if they persist on the self-interest hypothesis and rule out any social preferences. Nevertheless, even if they deal with inequity aversion, fairness, altruism, reciprocity and so forth, there is no clear outcome observable as the decision whether a fair or selfish agent is chosen depends on the assumptions made in their models.

In case of individual work without any interaction the fair agent of the model “Fairness in Delegated Bargaining”⁹⁸ will be preferred as he is also concerned about the principal’s payoff. Compared to the model “Moral hazard and other-regarding preferences”⁹⁹ where the selfish agent is preferred as it is more costly for the principal to induce the fair agent to exert more effort.

In case of interaction with an internal third party generally the fair agent is chosen except if the agent only cares for the third party. Compared to the model of Itoh¹⁰⁰, if the cost of effort is included the fair agent is chosen. Otherwise rather the selfish agent is preferred.

Interestingly would be if the results even hold under a repeated game. Perhaps it is possible if the fair agent cares too less for the third party, the principal in the next bargaining will choose the selfish agent.

⁹⁸ See Lammers, F. (2010), page 169ff

⁹⁹ See Itoh, H. (2004), page 18ff

¹⁰⁰ See Itoh, H (2004), page 18ff

8.2 German

Wie man sieht scheitern Ökonomen, wenn sie soziale Präferenzen ausschließen, da sie in der Realität sehr wohl eine Rolle spielen. Doch obwohl viele Ungleichheitsaversion, Fairness, Altruismus und dergleichen in ihren Modellen einbeziehen, ist kein eindeutiges Ergebnis, ob ein fairer oder eigennütziger Agent bevorzugt wird, sichtbar. Das Ergebnis der Entscheidung hängt stark von den getroffenen Annahmen in den Modellen ab.

Im Falle von individueller Arbeit ohne jegliche Interaktion wird im Modell „Fairness in Delegated Bargaining“¹⁰¹ klar der faire Agent bevorzugt, da er neben seiner Auszahlung auch die Auszahlung des Prinzipals in Erwägung zieht. Im Vergleich zum Modell „Moral hazard and other-regarding preferences“ von Herrn Itoh¹⁰², wird der egoistische Agent bevorzugt, da es für den Prinzipal teurer ist den fairen Agenten dazu zu bewegen mehr Anstrengungen zu zeigen.

Im Falle einer Interaktion mit einem internen Dritten wählt der Prinzipal grundsätzlich den fairen Agent, es sei denn dieser kümmert sich nur um den Dritten und keineswegs um den Prinzipal. Im Vergleich zum Modell von Itoh, wird der faire Agent bevorzugt, wenn die Kosten für die Anstrengungen im Modell berücksichtigt werden. Sonst wird eher der egoistische Agent bevorzugt.

Interessant wäre, ob die Ergebnisse auch im Fall von wiederholten Spielen halten. Hier wäre es eventuell möglich, falls sich der faire Agent zu wenig um die dritte Partei kümmert, der Prinzipal in den nächsten Tarifverhandlungen den egoistischen Agenten wählt?

¹⁰¹ See Lammers, F. (2010), page 169ff

¹⁰² See Itoh, H. (2004), page 18ff

Bibliography

Akerlof, G.A., Kranton, R.E. (2005): Identity and the Economics of Organizations, *Journal of Economic Perspectives*, Volume 19, Number 1, 9-32

Andreoni, J. (1989): Giving with impure altruism: Applications to Charity and Ricardian equivalence, *Journal of Political Economy*, Volume 97, Number 6, 1447–1458

Andreoni, J., Miller, J. (2002): Giving according to GARP: An Experimental Test of the Consistency of Preferences for Altruism, *Econometrica*, Volume 70, Number 2, 737-753

Aoki, M. (1984): *The Co-operative Game Theory of the Firm*, Oxford University Press, New York

Binmore, K., Rubinstein, A., Wolinsky, A. (1986): The Nash bargaining solution in economic modelling, *RAND Journal of Economics*, Volume 17, Number 2, 176-189

Bloch, F., Gomes, A. (2006): Contracting with Externalities and Outside Options, *Journal of Economic Theory*, Volume 127, Issue 1, 172-201

Bolton, G.E., Ockenfels, A. (2000): ERC: A Theory of Equity, Reciprocity, and Competition, *American Economic Review*, Volume 90, 166-193

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

Charness, G., Rabin, M. (2000): Social preferences: Some simple tests and a new model, mimeo, University of California at Berkeley, 1-71

Charness, G., Rabin, M. (2002): Understanding Social Preferences with Simple Tests, *The Quarterly Journal of Economics*, Volume 117, 817-869

Che, Y.K., Yoo, S.W. (2001): Optimal Incentives for Teams, *American Economic Review*, Volume 91, 817-869

Cox, J.C., Friedman, D., Gjerstad, S. (2009): A tractable Model of Reciprocity and Fairness, *Games and Economic Behavior*, Volume 59, 17-45

Dewatripont, M. (1988): Commitment through Renegotiation-Proof Contracts with Third Parties, *Review of Economic Studies*, Volume 55, Number 3, 377-390

Drago, R., Garvey, G.T. (1994): Incentives for Helping on the Job: Theory and Evidence, *Journal of Labor Economics*, Volume 16, 1-25

Dufwenberg, M., Kirchsteiner, G. (2004): A Theory of Sequential Reciprocity, *Games and Economic Behavior*, Volume 47, 268-298

Falk, A., Fehr, E., Fischbacher, U. (2000a): Informal sanctions, Institute for Empirical Research in Economics, University of Zurich, Working Paper No. 59, 1-38

Falk, A., Fischbacher, U. (2006): A Theory of Reciprocity, *Games and Economic Behavior*, Volume 54, 293-315

Fehr, E., Gächter, S. (1998): Reciprocity and Economics: The economic Implication of Homo Reciprocans, *European Economic Review*, Volume 42, 845-859

Fehr, E., Schmidt, K. (1999): A theory of Fairness, Competition and Cooperation, *The Quarterly Journal of Economics*, Volume 114, 817-868

Fehr, E., Schmidt, K. (2002): Theories of Fairness and Reciprocity – Evidence and Economic Applications, Working Paper No. 75, Working Paper Series ISSN 1424-0459, University of Zurich

Fehr, E., Fischbacher, U. (2002): Why social preferences matter: The impact of non-selfish motives on competition, cooperation and incentives, *The Economic Journal*, Volume 112, 1-33

Holmström, B. (1982): Moral Hazard in Teams, The Bell Journal of Economics, Volume 13, Number 2, 324-340

Itoh, H. (2004): Moral Hazard and other-regarding Preferences, The Japanese Economic Review, Volume 55, 18-45

Katz, M.L. (1991): Game-playing Agents: Unobservable Contracts as Precommitments, RAND Journal of Economics, Volume 22, 307-328

Holmstrom, B. (1982): Moral Hazard in Teams, Bell Journal of Economics, Volume 13, 324-340

Levin, D.K. (1998): Modeling Altruism and Spitefulness in Experiments, Review of Economic Dynamics, Volume 1, 593-622

Lammers, F. (2007): Fairness in Delegated Bargaining, Discussion Paper Series in Economics and Management, Discussion Paper No. 07-17, 1-34

Lammers, F. (2010): Fairness in Delegated Bargaining, Journal of Economics and Management Strategy, Volume 19, Number 1, 169-183

Mayer, B., Pfeiffer, T. (2004): Prinzipien der Anreizgestaltung bei Risikoaversion und sozialen Präferenzen, Zeitschrift für Betriebswirtschaft, 74. Jg. 2004, Number 9, Betriebswirtschaftlicher Verlag Dr. Th. Gabler GmbH, Wiesbaden, 1047-1075

Neilson, W.S., Stowe, J. (2003): Incentive Pay for other-regarding Workers, Duke University, Mimeo

Rabin, M. (1993): Incorporating Fairness in Game Theory and Economics, The American Economic Review, Volume 83, Number 5, 1281-1302

Rubinstein, A. (1982): Perfect Equilibrium in a Bargaining Model, Econometrica, Volume 50, 97-109

Sappington, D. (1983): Limited Liability Contracts between Principal and Agent, Journal of Economic Theory, Volume 29, 1-21

Simon, W. (2006): Persönlichkeitsmodelle und Persönlichkeitstests, GABAL Verlag GmbH, Offenbach

Curriculum vitae

Persönliche Daten

Name: Nicole FROMVALD, Bakk.
Lavaterstraße 3/2/14
1220 Wien

E-Mail-Adresse: nicole.fromvald@gmail.com

Geburtsdatum: 23. Jänner 1986

Ausbildung

2009 - 2012 Masterstudium: Betriebswirtschaft an der Universität Wien
– BWZ in 1210 Wien
KFK: Controlling und Industrial Management

Thema der Magisterarbeit: „Social Preferences in Delegated Bargaining“

2006 – 2009 Bakkalaureatsstudium: Betriebswirtschaft an der Universität Wien – BWZ in 1210 Wien
Vertiefung: Organisations- und Personalmanagement, Produktions-, Supply Chain und Strategisches Innovationsmanagement und Production Analysis

Thema der 1. Bakkalaureatsarbeit: „Motivation von Volunteers“ (in Organisations- und Produktionsmanagement)

Thema der 2. Bakkalaureatsarbeit: „Unternehmensanalyse der Schoeller-Bleckmann AG“ (in Produktionsmanagement)

2000 – 2006 HBLA – Höhere Bundeslehranstalt für Tourismus und wirtschaftliche Berufe in 1210 Wien
Zweig: Internationale Kommunikation in der Wirtschaft

1996 – 2000 Fremdsprachen Hauptschule in 1100 Wien
Sprache: Italienisch

Erfahrungen

Seit Juni 2011 Junior Controllerin bei REWE International AG (Penny)

November 2008 – Juni 2011 Geringfügige Beschäftigung bei REWE International AG (zentrales HR-Management)

August 2006 Praktikum bei Mobilfunkbetreiber A1 (Marketingabteilung für Business-Kunden)