

Diskussionspapier des
Instituts für Organisationsökonomik

1/2021

Cheating Alone and in Teams

Alexander Dilger

Discussion Paper of the
Institute for Organisational Economics

**Diskussionspapier des
Instituts für Organisationsökonomik
1/2021**

Januar 2021

ISSN 2191-2475

Cheating Alone and in Teams

Alexander Dilger

Abstract

There is a reward for a project that can be increased through ability, effort, and cheating. This is analysed for one agent and a team of two. As an extension, a preference for honesty is added, which can prevent cheating but not without limit and not so easily in the team context.

JEL Codes: D82, K42, Z20, Z22

Täuschen allein und in Teams

Zusammenfassung

Es gibt für ein Projekt eine Belohnung, die durch Fähigkeit, Anstrengung und Betrug gesteigert werden kann. Dies wird für einen Agenten und ein Zweierteam analysiert. Als Erweiterung wird eine Präferenz für Ehrlichkeit hinzugefügt, die Betrug verhindern kann, aber nicht unbegrenzt und nicht so leicht im Teamkontext.

Im Internet unter:

http://www.wiwi.uni-muenster.de/io/forschen/downloads/DP-IO_01_2021

Westfälische Wilhelms-Universität Münster
Institut für Organisationsökonomik
Scharnhorststraße 100
D-48151 Münster

Tel: +49-251/83-24303 (Sekretariat)
E-Mail: io@uni-muenster.de
Internet: www.wiwi.uni-muenster.de/io

Cheating Alone and in Teams

1. Introduction

Dilger (2017) analyses “Doping in Teams” that compete in a Tullock (1980) contest. This paper takes another look at cheating in general, of which doping in sports is just one special case. Other examples are fraudulent accounting, inflated bills by craftsmen or doctors, and also data falsification and plagiarism by scientists (cf. Dilger/Frick/Tolsdorf 2007: 604). In all these cases, the reward R of a project depends on ability a , effort e and, possibly, deceit d . Section 2 analyses the decision by one person whether and, if so, how much to cheat. Section 3 looks at the same decisions in a team of two. Section 4 considers an extension of the model with a preference not to cheat by oneself. Section 5 concludes.

2. Cheating Alone

There is a project with reward R depending on ability a , effort e and, possibly, deceit d according to (1):

$$(1) \quad R = ae + d.$$

That means that the reward depends on both ability and effort. They can substitute each other, although not completely, and they amplify each other. However, a cheater does not need (much) talent or effort to write more on the bill, to dope or to invent better data and results. Absolutely honest people do not cheat regardless, but opportunistic agents have to be limited by possible punishment P . They are punished with probability p depending on d :

$$(2) \quad p(d) = d^2.$$

This is only valid for $d \leq 1$, otherwise $p = 1$ because a probability cannot be larger than 1. However, here and in the following it is assumed that either d cannot become larger than 1 or that this is not worthwhile because P is always larger than R . Otherwise, open deceit that is detected for sure would be optimal. However, it would not be really deceit as everybody knows about it. It would be like legalising doping or plagiarism.

If the disutility of effort has a quadratic form (e^2), too, then the utility of the project is:

$$(3) \quad U = R - e^2 - d^2P = ae + d - e^2 - d^2P.$$

From this the optimal effort level e^* can be derived:

$$(4) \quad \delta U / \delta e = a - 2e^* = 0 \Rightarrow e^* = a/2.$$

This means more able agents will put in more effort, too, and get higher rewards for more successful projects.

The optimal degree of cheating for an opportunistic utility maximizer is:

$$(5) \quad \delta U / \delta d = 1 - 2d^*P = 0 \Rightarrow d^* = 1/(2P).$$

The degree of cheating is independent of ability and effort, at least in this simple model where they are separable components. The cheating depends only on the detection probability and the size of punishment in case of detection. Higher sanctions lead to less cheating. Nevertheless, there will always be some cheating by opportunists. How this is may be changed by a preference for honesty is analysed in Section 4 below.

3. Cheating in Teams

There is a team with two members, 1 and 2. They get the reward R for the project together. It depends on their abilities a_1 and a_2 , their efforts e_1 and e_2 and their deceptions d_1 and d_2 :

$$(6) \quad R = a_1e_1 + a_2e_2 + d_1 + d_2.$$

Once again, there is disutility of effort in a quadratic form and possible punishment with a detection probability also with a quadratic form depending on the sum of deceit. The punishments P_1 and P_2 can be different, but not so different that one of the team members does not want to participate in the team because the other could make his or her utility negative by cheating a lot with a low punishment while the own punishment is high.

The team members have these utility functions if team member 1 gets a share α of R and 2 gets $(1 - \alpha)$:

$$(7) \quad U_1 = \alpha R - e_1^2 - (d_1 + d_2)^2 P_1 = \alpha(a_1e_1 + a_2e_2 + d_1 + d_2) - e_1^2 - (d_1 + d_2)^2 P_1,$$

$$(8) \quad U_2 = (1 - \alpha)R - e_2^2 - (d_1 + d_2)^2 P_2 = (1 - \alpha)(a_1e_1 + a_2e_2 + d_1 + d_2) - e_2^2 - (d_1 + d_2)^2 P_2.$$

As before, from this the optimal effort levels e_1^* and e_2^* can be derived:

$$(9) \quad \delta U_1 / \delta e_1 = \alpha a_1 - 2e_1^* = 0 \Rightarrow e_1^* = \alpha a_1 / 2,$$

$$(10) \quad \delta U_2 / \delta e_2 = (1 - \alpha)a_2 - 2e_2^* = 0 \Rightarrow e_2^* = (1 - \alpha)a_2 / 2,$$

Comparing (9) and (10) with (4) shows that the effort incentives are reduced because of sharing the reward. With regard to cheating, there are these first-order conditions:

$$(11) \quad \delta U_1 / \delta d_1 = \alpha - 2d_1^*P_1 - 2d_2P_1 = 0 \Rightarrow d_1^* = \alpha / (2P_1) - d_2 \Rightarrow d_1^* = \alpha / (2P_1) - d_2^*,$$

$$(12) \quad \delta U_2 / \delta d_2 = (1 - \alpha) - 2d_1P_2 - 2d_2^*P_2 = 0 \Rightarrow d_2^* = (1 - \alpha) / (2P_2) - d_1 \Rightarrow d_2^* = (1 - \alpha) / (2P_2) - d_1^*.$$

Putting (12) into (11) to solve for d_1^* shows a problem:

$$(13) \quad d_1^* = \alpha / (2P_1) - [(1 - \alpha) / (2P_2) - d_1^*] \Rightarrow d_1^* = \alpha / (2P_1) - (1 - \alpha) / (2P_2) + d_1^* \Rightarrow (1 - \alpha) / (2P_2) = \alpha / (2P_1).$$

d_1^* stands on both sides of the equation, cancelling each other out. The remaining equation is not fulfilled for most parameter constellations. This means there are no inner solutions and d_1^* and/or d_2^* has to be 0 (or 1 but this has been ruled out before). If $(1 - \alpha) / (2P_2) > \alpha / (2P_1)$, then

$$(14) \quad d_1^* = 0 \text{ and } d_2^* = (1 - \alpha) / (2P_2).$$

If $(1 - \alpha) / (2P_2) < \alpha / (2P_1)$, then

$$(15) \quad d_1^* = \alpha / (2P_1) \text{ and } d_2^* = 0.$$

In both cases the team member does all the cheating who prefers it more because of a relatively higher share of the reward compared to his penalty if caught. The other team member would prefer less cheating and thus abstains from it. If the remaining equation in (13) is fulfilled, that is $(1 - \alpha) / (2P_2) = \alpha / (2P_1)$, there are infinite many pairs of d_1^* and d_2^* compatible with (11) and (12). In this case a higher d_1 just substitutes for a lower d_2 and vice versa. The joint cheating of the team is

$$(16) \quad d_1^* + d_2^* = (1 - \alpha) / (2P_2) = \alpha / (2P_1).$$

Comparing (14) to (16) with (5) shows that the incentives to cheat are reduced in the team context. However, the main reason for this is that both team members are punished. If their joint punishment is not higher for someone doing the project alone, that is $P_1 + P_2 = P$, then the level of deceit is comparable or for symmetric team members with $\alpha = 1/2$ and $P_1 = P_2$ it is even identical.

4. Preference for Honesty

There are many possible extension of the simple model used in Section 2 and Section 3. Here a preference for honesty is included and analysed. Honest agents could be modelled as lacking the option of deceit, in which case there would be a simple model without cheating, or as deriving disutility from cheating. In the second case, the disutility does not have to be infinite such that the honesty is not absolute. This seems to be true for most human beings who are neither totally opportunistic nor absolutely honest. If the disutility of cheating is $-D$ in a project done alone regardless of the amount of cheating, such an honesty preferring opportunist will only cheat under the condition that the advantage of cheating (d^*) is larger than its disadvantages ($D + d^{*2}P$) in her utility function (3) complimented with $-D$:

$$(17) \quad d^* - D - d^{*2}P > 0.$$

If (17) is the case, she will cheat according to (5). If not, she will not cheat at all. Thus, the punishment needed to prevent cheating by such an agent with a preference for honesty is

$$(18) \quad P \geq 1/d^* - D/d^{*2} = 2P - 4P^2D = 1/(4D).$$

Extending this to the team setting, some different cases have to be distinguished. If only one team member has this preferences for honesty regarding her own behaviour and she is not the one deceiving anyway, then the other team member will cheat as before. In the special case of joint cheating of (16), the member with this preferences for honesty does not cheat, the other one does all the cheating and the amount of cheating overall is unchanged.

Really interesting is the case that the team member with this preference for her own honesty is the one who would be cheating alone without this preference. It is assumed that she is team member 1 (the other case that she is team member 2 is equivalent). Then her deceit is either $d_1^* = \alpha/(2P_1)$ according to (15) or she does not deceit at all. If team member 2 would not cheat in any case ($d_2^* = 0$), the condition for her cheating would be equivalent to (17):

$$(19) \quad \alpha d_1^* - D_1 - d_1^{*2}P_1 > 0.$$

However, an unlimitedly opportunistic team member 2 would cheat if team member 1 does not. In this case, he would choose $d_2^* = (1 - \alpha)/(2P_2)$ according to (14). Team member 1 knows this and compares her utility as shown in (7) with either $d_1^* = \alpha/(2P_1)$, $-D_1$ and $d_2^* = 0$ or $d_1^* = 0$ and $d_2^* = (1 - \alpha)/(2P_2)$:

$$(20) \quad \alpha d_1^* - D_1 - d_1^{*2} P_1 = \alpha^2 / (4P_1) - D_1 > \alpha d_2^* - d_2^{*2} P_1 = \alpha(1 - \alpha) / (2P_2) - (1 - \alpha)^2 P_1 / (2P_2)^2 \\ \Rightarrow D_1 < \alpha^2 / (4P_1) - \alpha / (2P_2) + \alpha^2 / (2P_2) + (1 - \alpha)^2 P_1 / (2P_2)^2$$

Depending on the parameter values, it is possible (if $\alpha d_2^* - d_2^{*2} P_1 > \alpha d_1^* - D_1 - d_1^{*2} P_1 > 0$) that team member 1 decides to let the other one cheat to save D_1 although she would have cheated herself if he would not. If the disutility of team member 1 arises from any cheating in the project, not only that by herself, she would cheat always in this situation like having no disutility because without cheating her the other team member would cheat anyway.

If both team members have the same scruples to cheat by themselves, $D_1 = D_2$, condition (19) is binding and either team member 1 cheats or no one does. She could hope that the other team member does the cheating and bears its costs but her incentive for cheating is higher, which is a disadvantage in this case, while cheating of both is inefficient. The same reasoning applies for $D_1 < D_2$, whereas $D_1 > D_2$ could bring about one more condition for team member 2 comparable to (19), whether he is ready to cheat at all if team member 1 abstains from it:

$$(19) \quad (1 - \alpha) d_2^* - D_2 - d_2^{*2} P_2 > 0.$$

5. Conclusion

In the basic model of this paper, there is always some cheating even with very high punishment because the detection probability falls faster than the benefits of an increasingly small deceit. In the team context, one of the team members concentrates on cheating and the other one abstains because her net reward for cheating including possible punishment is lower. In case of a finite preference for honesty there is a threshold below which there is no cheating and above the same cheating occurs as without this preference. In the team context it is relevant who has this preference. There is no difference if it is the team member that is not cheating anyway. If it is the other team member, she could abstain from cheating but then the otherwise abstaining team member starts to cheat at the lower level preferred by him. Because she knows this, she could abstain even if she would have not in case of an honest companion who would not take over the cheating from her.

Literature

Dilger, Alexander (2017): “Doping in Teams: A Simple Decision Theoretic Model”, Discussion Paper of the Institute for Organisational Economics 6/2017, Münster.

- Dilger, Alexander/Frick, Bernd/Tolsdorf, Frank (2007): "Are Athletes Doped? Some Theoretical Arguments and Empirical Evidence", *Contemporary Economic Policy* 25(4), 604-615.
- Tullock, Gordon (1980): "Efficient Rent Seeking", in: James M. Buchanan/Robert D. Tollison/ & G. Tullock (Eds.): "Toward a Theory of the Rent-seeking Society", Texas A&M University Press, College Station (TX), pp. 97-112.

Diskussionspapiere des Instituts für Organisationsökonomik

Seit Institutsgründung im Oktober 2010 erscheint monatlich ein Diskussionspapier. Im Folgenden werden die letzten zwölf aufgeführt. Eine vollständige Liste mit Downloadmöglichkeit findet sich unter <http://www.wiwi.uni-muenster.de/io/de/forschen/diskussionspapiere>.

- DP-IO 1/2021** Cheating Alone and in Teams
Alexander Dilger
Januar 2021
- DP-IO 12/2020** Liberale Corona-Politik
Alexander Dilger
Dezember 2020
- DP-IO 11/2020** Abfindungen für Vorstandsmitglieder ohne und mit Beschränkungen
Alexander Dilger
November 2020
- DP-IO 10/2020** 10. Jahresbericht des Instituts für Organisationsökonomik
Alexander Dilger/Lars Vischer
Oktober 2020
- DP-IO 9/2020** Stellungnahme zur Änderung des Brennstoffemissionshandelsgesetzes
Alexander Dilger
September 2020
- DP-IO 8/2020** Sind Klausuren überflüssig?
Zum Zusammenhang zwischen PISA-Ergebnissen, Rechtschreibung und Noten
Alexander Dilger
August 2020
- DP-IO 7/2020** No Home Bias in Ghost Games
Alexander Dilger/Lars Vischer
Juli 2020
- DP-IO 6/2020** The Advances of Community Cloud Computing in the Business-to-Business-Buying Process
Nicolas Henn/Todor S. Lohwasser
Juni 2020
- DP-IO 5/2020** Wirtschaftsethische Überlegungen zum Klimawandel
Alexander Dilger
Mai 2020
- DP-IO 4/2020** Meta-Analyzing the Relative Performance of Venture Capital-Backed Firms
Todor S. Lohwasser
April 2020
- DP-IO 3/2020** From Signalling to Endorsement
The Valorisation of Fledgling Digital Ventures
Milan Frederik Klus
März 2020
- DP-IO 2/2020** Internet-Publikationen gehört die Zukunft
Alexander Dilger
Februar 2020



Herausgeber:
Prof. Dr. Alexander Dilger
Westfälische Wilhelms-Universität Münster
Institut für Organisationsökonomik
Scharnhorststr. 100
D-48151 Münster

Tel: +49-251/83-24303

Fax: +49-251/83-28429

www.wiwi.uni-muenster.de/io

