

Netzer, Nick; Scheuer, Florian

Working Paper

Competitive Markets without Commitment

Working Paper, No. 0814

Provided in Cooperation with:

Socioeconomic Institute (SOI), University of Zurich

Suggested Citation: Netzer, Nick; Scheuer, Florian (2008) : Competitive Markets without Commitment, Working Paper, No. 0814, University of Zurich, Socioeconomic Institute, Zurich

This Version is available at:

<https://hdl.handle.net/10419/76179>

Standard-Nutzungsbedingungen:

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

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

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

Terms of use:

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

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

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



University of Zurich

Socioeconomic Institute
Sozialökonomisches Institut

Working Paper No. 0814

Competitive Markets without Commitment

Nick Netzer, Florian Scheuer

November 2008

Socioeconomic Institute
University of Zurich

Working Paper No. 0814

Competitive Markets without Commitment

November 2008

Author's address: Nick Netzer
E-mail: nick.netzer@soi.uzh.ch

Florian Scheuer
E-mail: scheuer@mit.edu

Publisher Sozialökonomisches Institut
Bibliothek (Working Paper)
Rämistrasse 71
CH-8006 Zürich
Phone: +41-44-634 21 37
Fax: +41-44-634 49 82
URL: www.soi.uzh.ch
E-mail: soilib@soi.uzh.ch

Competitive Markets without Commitment*

Nick Netzer
University of Zurich

Florian Scheuer
MIT

November 2008

Abstract

In the presence of a time-inconsistency problem with optimal agency contracts, we show that competitive markets implement allocations that Pareto dominate those achieved by a benevolent planner, they induce strictly more effort, and they sometimes make the commitment problem disappear entirely. In particular, we analyze a model with moral hazard and two-sided lack of commitment. After agents have chosen a hidden effort and the need to provide incentives has vanished, firms can modify their contracts and agents can switch firms. As long as the ex-post market outcome satisfies a weak notion of competitiveness and sufficiently separates individuals who choose different effort levels, the market allocation is Pareto superior to a social planner's allocation. We construct a specific market game that naturally generates robust equilibria with these properties. In addition, we show that equilibrium contracts without commitment are identical to those with full commitment if the latter involve no cross-subsidization between individuals who choose different effort levels.

*Email addresses: nick.netzer@soi.uzh.ch and scheuer@mit.edu. We are grateful to Daron Acemoglu, Carlos Alós-Ferrer, Helmut Bester, Peter Diamond, Dennis Gaertner, Mike Golosov, Jon Gruber, Bard Harstad, Casey Rothschild, Armin Schmutzler, Jean Tirole, Iván Werning, and seminar participants at FU and HU Berlin, MIT, Northwestern University, and the University of Zurich for valuable suggestions. All errors are our own.

1 Introduction

Optimal contracts in the presence of moral hazard, where a risk-averse agent is able to affect the probability distribution over output by choosing some hidden action, reflect a tradeoff between providing incentives and insurance. However, such contracts are subject to a fundamental time-inconsistency problem: Whereas underinsurance is typically optimal *ex ante* so that the agent has incentives to exert effort, it becomes suboptimal once effort has been chosen. Since the need to provide incentives has vanished, a risk-neutral principal would find it optimal to provide the agent with full insurance after effort choice has been made, but before output is fully realized. The agent, anticipating this, then would have incentives to exert the least costly effort level.

The mechanism to deal with commitment problems in this and other related settings that has received most attention is *reputation*.¹ Repeated interaction between the principal and the agent allows to avoid the outcome with lowest effort and full insurance based on the credible threat of punishments in future periods. In this paper, we examine how an alternative mechanism performs in the framework of the outlined time-inconsistency problem: *competitive markets*.² Notably, we consider a model where contracts are offered by many competing principals in a market, and there is two-sided lack of commitment: Principals are unable to commit to contracts before agents choose their hidden effort, and agents are free to choose different contracts or principals after they have taken their effort decision.³ We demonstrate that this form of competition without commitment is able to deal with the time-inconsistency problem very successfully, even without allowing for any reputational effects.

To evaluate the performance of competitive markets without commitment relative to other institutions, we start with considering a government faced with the same informational and commitment constraints. In particular, it offers incentive contracts to a population of agents who differ in their privately known disutility of effort, and is free to change them after agents have chosen their effort and hence their *ex post* type, but before output is realized.

¹Zhao (2006) studies the above mentioned problem of moral hazard without commitment in an infinitely repeated principal agent model. Other applications, where reputational concerns can affect a commitment problem, include monetary policy (Kydland and Prescott (1977), Barro and Gordon (1983)) and capital taxation (Chari and Kehoe (1990), Phelan and Stacchetti (2001), Farhi and Werning (2008)).

²With the notable exception of some contributions to capital tax competition (e.g. Kehoe (1989) or Conconi, Perroni, and Riezman (2008)), the effect of competition on time-inconsistency problems has received little attention. The tax competition literature is concerned with the optimal degree of cooperation between countries, facing a trade-off between disciplining effects of non-cooperative behavior and an adverse race-to-the-bottom.

³In contrast, Phelan (1995) and Krueger and Uhlig (2006) study a competitive market structure with one-sided commitment, where, in the context of optimal risk sharing without an *ex ante* moral hazard problem, financial intermediaries are committed to contracts but agents are free to switch.

We first show that, with a utilitarian social planner, the unique equilibrium is such that no agent provides effort and everybody is fully insured.⁴ The same holds whenever the planner attaches overproportional Pareto-weights to agents with a high effort cost compared to their population share, which may be a particularly realistic case in many applications where the government is driven by redistributive concerns. This is because the social planner in this case still has incentives to provide full insurance (or even overinsurance) ex post, which eliminates ex ante incentives.

We then turn to the analysis of competitive markets, where a large set of risk-neutral principals (firms) offer contracts to screen agents. Firms can offer new and modify their old contracts, and agents can move to other firms once they have chosen effort.⁵ Then, firms take the agents' effort decision and the composition of the population of agents as given in the ex post stage. We first show that, whenever the outcome of an ex post market satisfies a weak notion of competitiveness (minimal contestability) and it sufficiently separates agents who have taken different effort choices, it Pareto-dominates the government outcome. The reason is that, whereas incentives for effort completely break down if contracts are provided by a utilitarian government, this is not the case with competitive markets even in our one-shot economy without reputation effects. Importantly, while this result depends on the assumption that markets are competitive, it does not rely on a detailed specification of a particular market game. However, our results do not apply to the comparison between a monopolistic private firm and a government: competition is crucial.

We next construct a game-theoretic foundation for our analysis of competitive markets that generates allocations with the desired properties as sequential equilibrium outcomes. The extensive form is such that, first, agents take a hidden effort decision. Then, each firm forms a belief about agents' effort choices and offers a finite number of incentive contracts. Following the suggestion by Miyazaki (1977), we introduce a withdrawal phase to guarantee equilibrium existence: After observing the set of offered contracts, firms can decide to become inactive at a small cost. Finally, agents choose one of the remaining contracts, and the uncertainty is realized. The withdrawal phase yields equilibrium existence in the ex post market, but also equilibrium multiplicity.⁶ We show that the multiplicity problem can be solved by introducing a simple robustness requirement based on varying withdrawal costs,

⁴The mechanism underlying this result relates to what has been described as the Samaritan's dilemma (Buchanan 1975) or the problem of soft budget constraints (Kornai, Maskin, and Roland 2003). The same result, which predicts a complete breakdown of incentives, has been obtained by Boudway, Marceau, and Marchand (1996) in the setting of education and redistributive tax policy.

⁵Our analysis thus applies whenever a switching agent can take along its type to the new firm. This is plausible in many applications, as we will discuss below.

⁶A similar problem has been observed by Engers and Fernandez (1987) in a game with an infinite sequence of contract offers.

which yields Miyazaki-Wilson type contracts as the unique robust outcome in the ex post market with given effort choice. These contracts correspond to the optimal solution a social planner would implement who places welfare weight only on high effort types.

We then provide a characterization of robust sequential equilibrium outcomes in the complete market game without commitment, including the agents' optimal effort choice. We show that they are the solution to a fixed point problem, where the agents' ex ante effort decision must be optimal given the resulting ex post equilibrium contracts, and vice versa. We provide conditions for the existence and uniqueness of such equilibria. Moreover, we show that sequential equilibrium outcomes Pareto-dominate the allocation with a utilitarian planner from both an ex ante and ex post perspective, i.e. all effort cost types ex ante prefer the market allocation to that implemented by the planner, and the same holds for all ex post types.

We consider three main extensions of this result. First, we assume that, even when competitive markets are in place, the government may shut them down ex post with some exogenous probability, and implement the policy it considers optimal. Whenever the probability that the market outcome is destroyed ex post is strictly less than one (but possibly arbitrarily close to one), our Pareto comparison goes through. Intuitively, if agents anticipate that the market outcome has some chance of being implemented, there remain some ex ante incentives left, in contrast to the equilibrium with a government only.

Second, whereas the Pareto comparisons depend on the social planner being concerned about a utilitarian welfare criterion or one that makes redistribution towards high effort cost types desirable, we show that, for any distribution of Pareto-weights that the planner may use to evaluate welfare, competitive markets without commitment implement more effort than the social planner with the same commitment problem under general conditions. Intuitively, as our ex post market replicates a social planner who cares only about high effort types, it provides maximal incentives for effort ex ante.

Finally, we compare competitive markets without commitment to the full commitment benchmark. Here, firms are able to sign binding contracts with agents *before* the hidden effort choice. Thus, they take the effect of their contract offers on the ex post composition of their agent pool fully into account. We find that, with competitive markets, the commitment problem disappears for an entire class of economies: Whenever an equilibrium with full commitment involves no cross-subsidization between ex post types, i.e. between agents who have chosen different effort levels, it is also an equilibrium if there is no commitment. In that case, competitive markets fully solve the commitment problem. Otherwise, lack of commitment leads to less effort and more equilibrium insurance than in the full commitment case. However, even in this case, competitive markets are able to provide effort incentives

for a non-zero mass of agents.

In terms of economic settings to which our model applies, consider for instance an individual's education decision, as in Boadway, Marceau, and Marchand (1996) and Konrad (2001), and subsequent labor markets. Significant parts of education are private information and are typically completed before binding contracts with employers are signed. Even if there are contracts, as in the case of executive education, agents cannot be prevented from moving to other employers some time after their education has been completed, and employers are able to modify employment contracts or eventually lay off employees. We characterize the equilibrium level of education and the form of employment contracts in such a setting when there is competition between employers. Another example for our general model are insurance markets with ex ante moral hazard. Here, agents can influence their risk of a negative outcome, such as a damage, illness or unemployment, by choosing some hidden effort. Once preventive effort has been chosen, but before the risk is realized, competitive insurance companies can modify contractual terms and customers can switch insurers.

The paper most closely related to ours is the seminal contribution by Fudenberg and Tirole (1990). They observe the same time-inconsistency problem that we analyze here in a principal-agent economy where a monopolistic risk-neutral principal designs the optimal contract for a risk-averse agent but cannot commit not to renegotiate it in the interim stage after effort choice.⁷ When contracts are provided by competing principals as opposed to a monopolist, the commitment problem becomes of a different nature. Notably, the issue of two-sided lack of commitment arises, since agents may switch firms after choosing their hidden effort. However, although this additional commitment problem may be expected to lead to even worse equilibrium outcomes, it turns out that, with competitive markets, the equilibrium cross-subsidization between ex post effort types is crucial in determining whether there is a commitment problem at all. Moreover, Fudenberg and Tirole (1990) do not compare the performance of different institutions in view of the commitment problem.

Asheim and Nilssen (1996) study renegotiation in competitive insurance markets, but the time-inconsistency problem that they consider is different from the moral hazard problem analyzed here. Risk types are exogenous in their model, and insurance firms renegotiate the contracts with their customers after they have observed initial contract choices, making use of the information revealed through choice.⁸ Yet another type of commitment problem occurs

⁷Since our model includes heterogeneity in effort costs, the equilibria that we consider do not involve agents randomizing between effort levels, as opposed to the equilibria in Fudenberg and Tirole (1990). In their section 3, Fudenberg and Tirole (1990) outline the possibility of purification through effort cost heterogeneity.

⁸Bester and Strausz (2001) investigate the validity of the revelation principle in a general mechanism design setting with a similar commitment problem.

when insurance companies can use individual loss histories to update their information on (exogenous) risk types. This problem has been investigated by Kunreuther and Pauly (1985), Dionne and Doherty (1994) and Nilssen (2000) for the case of competitive markets.⁹ Since there is no time-inconsistency problem related to moral hazard in these models, they are not able to make predictions about how the commitment problem affects ex ante effort incentives, which is what we are interested in here. Thus, we ignore these types of commitment problems in the present paper.

Comparing the efficiency of markets and governments in a setting without commitment, our paper shares a common goal with the contributions by Acemoglu, Golosov, and Tsyvinski (2008a, 2008b). However, their modelling of both markets and governments is very different from the approach taken here. The provision of insurance contracts by private firms in competitive markets is ruled out, and government policies are distorted by political economy constraints. Moreover, their equilibria crucially rely on reputational concerns in an infinitely repeated game. In contrast, we completely abstract from reputational effects, assume a benevolent government and consider markets where competitive firms can offer insurance contracts that are only restricted by informational and commitment constraints.

Finally, our results also complement a vast literature on public versus private provision of goods and services in an incomplete contracts world (see Shleifer (1998) for an overview). Whereas this literature focuses on how privatization affects the asymmetry of information or the production technology in a firm, we derive a clear advantage of competitive markets over a benevolent government without assuming any differences in the technological, informational or commitment constraints faced by these different institutions. In contrast to Schmidt (1996) and Bisin and Rampini (2006), where privatization or the creation of anonymous markets, respectively, is assumed to conceal information from the government and to act as a constraint on the set of feasible policies,¹⁰ we show that the establishment of competitive markets can be interpreted as the choice of a specific, effort prone welfare function.

The paper is structured as follows. Section 2 introduces our model economy. In section 3, we compare competitive equilibria without commitment to those achieved by a social planner with the same commitment problem. The analysis rests on a very general, axiomatic treatment of competitive market outcomes. We then provide a rigorous game-theoretic foundation for this procedure in section 4, where sequential equilibria of a specific market game are characterized. Section 5 contains the extensions, notably our comparison of competitive

⁹All these types of commitment problems, caused by the flow of additional information after initial contract conclusion, are in line with the literature on the ratchet effect (see Freixas, Guesnerie, and Tirole (1985) and Dewatripont (1989)).

¹⁰Konrad (2001) also highlights that having less information can be beneficial for a government in the presence of a commitment problem.

markets without commitment to the full commitment case, and section 6 concludes.

2 The Model

To study the issues raised in the previous section formally, we consider the following model economy. There is a continuum of risk-averse agents, indexed by the set $[0, \infty)$. Agents are expected utility maximizers with a Bernoulli utility function $U(c)$, where c is consumption. $U(c)$ is twice continuously differentiable, with $U' > 0$ and $U'' < 0$. Following Fudenberg and Tirole (1990), we assume that both the domain and the range of U are given by $\mathbb{R}$, so that $\lim_{c \rightarrow -\infty} U(c) = -\infty$ and $\lim_{c \rightarrow \infty} U(c) = \infty$.¹¹ We also assume that the Inada condition $\lim_{c \rightarrow \infty} U'(c) = 0$ is satisfied. Let $\Phi(U)$ be the inverse function of U , which then satisfies $\Phi' > 0$, $\Phi'' > 0$, $\lim_{U \rightarrow -\infty} \Phi(U) = -\infty$, $\lim_{U \rightarrow \infty} \Phi(U) = \infty$ and $\lim_{U \rightarrow \infty} \Phi'(U) = 0$.

Each agent faces idiosyncratic risk with respect to the amount of the consumption good that she produces for a firm (or principal). The output can either be high, y_h , or low, y_l , with $y_l < y_h$. If the agent is a good type (g), the probability of the high output is p_g . Bad types (b) produce the high output with probability p_b , where $0 < p_b < p_g < 1$ holds.¹² An agent's type depends on the effort level $e \in \{\underline{e}, \bar{e}\}$ that she exerts. Upon choosing the high effort $\bar{e}$, the agent becomes a good type, whereas an agent who chooses the low effort $\underline{e}$ becomes a bad type. We assume that a law of large numbers applies to the continuum of random variables defined by the population facing idiosyncratic risk. That is, we assume that exactly the share p_g (p_b) of any set of good (bad) type agents with positive measure does eventually produce the high output.¹³

The agents' preferences are assumed to be separable between consumption and effort, so that overall utility is given by $U(c) - H(e)$, where $H(e)$ denotes effort cost. We normalize $H(\underline{e})$ to zero. Agents differ in their disutility of effort $H(\bar{e}) = d$, which is given by their index $d \in [0, \infty)$. The composition of the population is described by a continuous distribution function G , defined on $\mathbb{R}$ with $G(d) = 0$ for all $d \leq 0$. We adopt the convention of extending

¹¹These assumptions can be relaxed to accommodate frequently used utility functions, such as those with constant absolute or relative risk aversion, but they simplify the proofs in the following.

¹²In an insurance market application, y_h represents each agent's endowment, and $y_h - y_l$ is a possible damage. The damage occurs with low probability $1 - p_g$ for good types, i.e. low risks, and with larger probability $1 - p_b$ for bad types, i.e. high risks. In other applications, where y_k , $k = l, h$, actually represents the amount of output produced in a firm, our partial equilibrium model could be embedded into a general equilibrium model under the assumption of price-taking behavior on the output market. See Schmidt (1997) for a model that analyzes managerial incentives under imperfect product market competition.

¹³While laws of large numbers for a continuum of random variables may fail due to technical complications (see Judd (1985)), they can be put back into force through a variety of approaches. These include the application of a weaker convergence criterion (Uhlig (1996)), the redefinition of the set indexing consumers (Green (1994)), or the derivation of individual risk from the desired aggregate level properties (Alós-Ferrer (2002)).

G to $G(\infty) = 1$. We also assume that G has an associated density g that satisfies $g(d) > 0$ for all $d \in [0, \infty)$.¹⁴ We assume throughout the paper that neither effort cost nor effort choice are observable to anyone besides the agent herself.

For our analysis, it is very convenient to operate in the utility space. In this space, a contract that a planner or a firm offers to an agent is a tuple (u_h, u_l) of consumption utilities that the agent obtains when producing the high and the low output, respectively. Let $\mathcal{I} = \{(u_h, u_l) \in \mathbb{R}^2 | u_h \geq u_l\}$ be the set of possible contracts.¹⁵ We denote the set of all finite, non-empty subsets of $\mathcal{I}$ by $\mathcal{Q}$. Finally, let $\mathcal{V} = \{(u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathbb{R}^4 | u_{b,h} \geq u_{b,l}, u_{g,h} \geq u_{g,l}\}$ be the set of quadruples representing pairs of contracts, one intended for bad types $(u_{b,h}, u_{b,l})$ and one for good types $(u_{g,h}, u_{g,l})$.

3 Markets Versus Governments

In this section, we are concerned with the comparison between a benevolent planner and competitive markets. We first demonstrate that a complete breakdown of incentives is the unique outcome if a social planner is in charge who maximizes some welfare function from a large class that includes, for example, the utilitarian case.¹⁶ Our subsequent analysis of competitive markets abstracts from details of a market game by proceeding axiomatically: We formulate plausible properties that outcomes of an ex post market should satisfy to capture a weak notion of competitiveness. Based on these properties only, we are able to prove a first result on the Pareto dominance of markets over a planner. In Section 4, we provide a game theoretic foundation for our approach by modelling an explicit game in which equilibrium outcomes satisfy the properties postulated in the following.

3.1 A Social Planner Without Commitment

We consider the following reduced timing:

Stage 1: Agents simultaneously choose their effort level.

Stage 2: The social planner announces a policy, i.e. a set of two contracts.

Stage 3: Agents simultaneously choose among the offered contracts.

¹⁴In fact, all our results go through with the somewhat weaker assumption that $g(d) > 0$ for $d = 0$ and for some sufficiently high d .

¹⁵Contracts with negative incentives $u_h - u_l < 0$ are not relevant for our analysis. Since agents are risk-averse, such contracts waste resources without being able to induce any effort.

¹⁶This result is a generalization of a similar one that Boadway, Marceau, and Marchand (1996) derive in the setting of education and redistributive taxation.

One could think of stage 1 being preceded by an additional stage where the planner announces an initial policy. Then, the agents choose an effort level and possibly one of the contracts that the planner has offered. If the planner is not committed to its initial announcement, however, she is free to change the policy ex post, after effort choice, and the initial offers become irrelevant.¹⁷ In an insurance application, the policy could be an optimally designed social insurance arrangement, such as a mandatory public health insurance where individuals can still choose between different levels of franchise.¹⁸ In an education and job market application, the policy describes the payment structure of jobs in the public sector, such as in schools or prisons (Hart, Shleifer, and Vishny 1997) or in other firms owned by the state (La Porta, Lopez-De-Silanes, and Shleifer 2002), or a redistributive tax policy (Boadway, Marceau, and Marchand 1996).

We solve the game backwards. First, for any given (unobservable) effort choice in stage 1 and policy announcement in stage 2 (consisting of two contracts), each ex post type $k \in \{g, b\}$ selects the best contract in stage 3, where we break ties in favor of the contract with larger insurance coverage, i.e. smaller difference $u_g - u_b$. Stage 3 can then be eliminated by subsuming this choice into the planner's payoff function. We can next derive the planner's optimal policy at stage 2 when effort choices have been made. Observe first, however, that optimal effort choices in stage 1 must be of a threshold type in any equilibrium, with a critical value $\hat{d} \in \mathbb{R}_0^+ \cup \{\infty\}$ such that

$$e(d) = \begin{cases} \bar{e} & \text{if } d < \hat{d}, \\ \underline{e} & \text{if } d \geq \hat{d}, \end{cases} \quad (1)$$

where $e(d)$ is the effort choice of an agent of type $d \in [0, \infty)$. Whenever an agent with effort cost d finds it optimal to choose the high effort, in anticipation of some final policy, the same holds for any agent of cost type $d' \leq d$. Thus, in any equilibrium, agents with small effort cost ($d < \hat{d}$) choose the high effort and those with high effort cost ($d \geq \hat{d}$) the low effort, and

¹⁷The irrelevance of initial offers rests on two assumptions. First, and as discussed in Section 1, we do not allow the planner to make use of information possibly revealed through the initial contract choice, as this might introduce an additional time-inconsistency problem different from the one we are interested in. This is not necessarily a restrictive assumption, however. A utilitarian planner, for example, does not benefit ex post from the additional information and would not use it anyway. Second, the initial announcement does not constitute binding reservation constraints to the planner, because there is no higher authority to enforce it. The two assumptions together imply that the planner is free to deviate from earlier policy announcements but cannot target specific individuals ex post. To allow for a reasonable comparison, we will impose analogous assumptions on firms in a market later.

¹⁸Observe that we allow the planner to offer two contracts, which enables her to design an unrestricted optimal policy. More than two contracts are indeed not necessary in the present setup with only two ex post types. Separability of effort cost implies that the ex ante heterogeneity cannot be used to screen the population ex post when effort is sunk.

the share of good types in the society becomes $G(\hat{d})$.¹⁹ Suppose then that the planner has formed a correct belief about $\hat{d}$ and uses a distribution of Pareto-weights $\Psi(d)$ for ex ante cost types $d \in [0, \infty)$ to evaluate social welfare in the economy. For example, the weights used by a utilitarian planner are the respective population shares, such that $\Psi(\hat{d}) = G(\hat{d})$. In general, varying welfare weights allow us to derive the whole ex post Pareto frontier.²⁰

Then, whenever $\hat{d} \in (0, \infty)$ so that both ex post types exist, the planner solves the following problem, which we refer to as program $SP(\hat{d})$:

$$\max_{(u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}} \Psi(\hat{d})[p_g u_{g,h} + (1 - p_g) u_{g,l}] + (1 - \Psi(\hat{d}))[p_b u_{b,h} + (1 - p_b) u_{b,l}] \quad (2)$$

subject to the constraints

$$p_g u_{g,h} + (1 - p_g) u_{g,l} \geq p_g u_{b,h} + (1 - p_g) u_{b,l}, \quad (3)$$

$$p_b u_{b,h} + (1 - p_b) u_{b,l} \geq p_b u_{g,h} + (1 - p_b) u_{g,l}, \quad (4)$$

$$G(\hat{d})[p_g \Phi(u_{g,h}) + (1 - p_g) \Phi(u_{g,l})] + (1 - G(\hat{d}))[p_b \Phi(u_{b,h}) + (1 - p_b) \Phi(u_{b,l})] \leq \mathbb{E}[\tilde{y}|\hat{d}]. \quad (5)$$

The planner maximizes a weighted average of the expected utilities of good and bad types subject to the two standard incentive constraints and the resource constraint. Here, $\tilde{y}$ is a Bernoulli random variable that takes the value y_h with probability $G(\hat{d})p_g + (1 - G(\hat{d}))p_b$ and y_l otherwise, and hence $\mathbb{E}[\tilde{y}|\hat{d}] = [G(\hat{d})p_g + (1 - G(\hat{d}))p_b]y_h + [1 - G(\hat{d})p_g - (1 - G(\hat{d}))p_b]y_l$ are the average resources per capita. The following lemma characterizes the solution of this problem:

Lemma 1. (i) For any given $\hat{d} \in (0, \infty)$, the program $SP(\hat{d})$ has a unique solution $V^{SP}(\hat{d}) = (u_{b,h}^{SP}(\hat{d}), u_{b,l}^{SP}(\hat{d}), u_{g,h}^{SP}(\hat{d}), u_{g,l}^{SP}(\hat{d}))$. It is such that $u_{b,h}^{SP}(\hat{d}) = u_{b,l}^{SP}(\hat{d})$ and constraints (4) and (5) are binding.

(ii) If $\Psi(\hat{d}) \leq G(\hat{d})$, $V^{SP}(\hat{d})$ satisfies $u_{b,h}^{SP}(\hat{d}) = u_{b,l}^{SP}(\hat{d}) = u_{g,h}^{SP}(\hat{d}) = u_{g,l}^{SP}(\hat{d}) = U(\mathbb{E}[\tilde{y}|\hat{d}]) \equiv u^{SP}(\hat{d})$.

Proof. See Appendix A.1. □

A social planner who puts a welfare weight on high effort cost agents that corresponds to or exceeds their population share always implements a pooling allocation with full insurance

¹⁹Also, since we assume that effort choice in stage 1 is unobservable, we do not need to derive the planner's optimal policy for effort choice profiles different from those threshold profiles, because potential deviations from an equilibrium candidate cannot be detected by the planner.

²⁰Also, the approach can be interpreted as the reduced form of a model in which the planner is concerned about effort and hence aggregate resources directly, or even takes into account individual effort costs in a non-separable manner. As long as her ex post policy is located on the ex post Pareto frontier, an appropriate weighting scheme Ψ can be found that induces the equivalent government behavior.

ex post.²¹ Incentive contracts with output-dependent utilities can only be optimal from such a planner's perspective if they increase effort investments. But if the effort choice has already been taken, the planner can always increase welfare by removing incentive components from the contracts and by pooling all agents into a single contract with output-independent utility. The assumption that a social planner puts weakly more weight on high effort cost types than on low cost types (relative to their population density) due to some redistributive concerns may be particularly realistic in many of the applications discussed above.

Clearly, the same holds if $\hat{d} \in \{0, \infty\}$, i.e. if all agents are either good types or bad types. In that case, the planner's problem $SP(\hat{d})$ prescribes the utility maximization of the unique ex post type, subject to a resource constraint. First, resources will clearly be exhausted in the solution. Second, convexity of Φ implies that the solution entails an output-independent payment. Hence, for $\hat{d} \in \{0, \infty\}$, we analogously define $V^{SP}(\hat{d}) = (u_{b,h}^{SP}(\hat{d}), u_{b,l}^{SP}(\hat{d}), u_{g,h}^{SP}(\hat{d}), u_{g,l}^{SP}(\hat{d}))$ by $u_{b,h}^{SP}(\hat{d}) = u_{b,l}^{SP}(\hat{d}) = u_{g,h}^{SP}(\hat{d}) = u_{g,l}^{SP}(\hat{d}) = \mathbb{E}[\tilde{y}|\hat{d}]$.²² Based on the planner's solution to $SP(\hat{d})$ for any $\hat{d} \in \mathbb{R}_0^+ \cup \{\infty\}$, we define an equilibrium with a social planner without commitment as follows.

Definition 1. *An equilibrium with a social planner without commitment (ESP) is a pair $(\hat{d}, V^{SP}(\hat{d}))$ where*

- (i) $V^{SP}(\hat{d})$ is a solution to $SP(\hat{d})$ and
- (ii) $\hat{d} = p_g u_{g,h}^{SP}(\hat{d}) + (1 - p_g) u_{g,l}^{SP}(\hat{d}) - p_b u_{b,h}^{SP}(\hat{d}) - (1 - p_b) u_{b,l}^{SP}(\hat{d})$.

The idea behind the fixed point condition (ii) is the following. Assume that agents' effort choices in stage 1 are given by a threshold $\hat{d}$ as described above, and the planner implements $V^{SP}(\hat{d})$ subsequently. Given that agents in stage 1 anticipate this, their actually optimal effort choice is described by the threshold

$$D^{SP}(\hat{d}) = p_g u_{g,h}^{SP}(\hat{d}) + (1 - p_g) u_{g,l}^{SP}(\hat{d}) - p_b u_{b,h}^{SP}(\hat{d}) - (1 - p_b) u_{b,l}^{SP}(\hat{d}).$$

This holds because each agent calculates her ex post utility from being a good type (choosing the good type's optimal contract) and the corresponding utility from being a bad type, and compares the difference to her effort cost d . The function $D^{SP}(\hat{d})$ therefore yields the indifferent cost type for any exogenously given $\hat{d}$, and no agent has an incentive to deviate if and only if the fixed point condition $\hat{d} = D^{SP}(\hat{d})$ and hence (ii) is satisfied.

²¹This is due our restriction that solutions must be in the set $\mathcal{V}$, which rules out overinsurance. Without that restriction, the government would ex post fully insure good types and overinsure bad types, leading to the same result that no incentives for effort can be provided.

²²We let $V^{SP}(0)$ and $V^{SP}(\infty)$ be elements of $\mathcal{V}$ for notational consistency, even though there is only one ex post type if $\hat{d} \in \{0, \infty\}$. One can still think of $(u_{k,h}^{SP}(\hat{d}), u_{k,l}^{SP}(\hat{d}))$ as the best contract for type $k \in \{g, b\}$ among those offered, even though only one type actually exists and the planner offers a single contract only.

Now assume that $\Psi \succeq_{FOSD} G$, i.e. $\Psi(\hat{d}) \leq G(\hat{d})$ for all $\hat{d}$. It then follows immediately from Lemma 1 that $D^{SP}(\hat{d}) = 0$ for all $\hat{d}$, which implies that $\hat{d}^{SP} = 0$ is trivially the only fixed point of $D^{SP}(\hat{d})$, and we obtain the following result:

Proposition 1. *If $\Psi \succeq_{FOSD} G$, then $(0, V^{SP}(0))$ is the unique ESP.*

If the government puts weakly more welfare weight on high cost types than is given by their population share, which is implied by the assumption $\Psi \succeq_{FOSD} G$, then it also puts overproportional weight on bad types ex post for any given $\hat{d}$. Continuation contracts are therefore such that all agents obtain full insurance, and no ex ante incentives can be sustained. We will return to the case where $\Psi \succeq_{FOSD} G$ is not satisfied in Section 5.2.

3.2 Competitive Markets Without Commitment

We now turn to the case where contracts are provided by competitive firms rather than a social planner, and both firms and agents are unable to commit to contracts before the hidden effort choice. That is, we consider a time structure with two-sided lack of commitment, where, after an initial phase of contract offers and agents' choices of contracts and effort, firms are free to alter their existing contracts and offer additional ones, while agents are free to abrogate their contract and choose a new one, possibly switching between firms.

It is again important to notice that we do not allow the firms' new contract offers or modifications to be conditioned on an agent's initial choice of contract. First, this allows us to isolate the effects of our time-inconsistency problem from those in Asheim and Nilssen (1996) and the literature on the ratchet effect (Freixas, Guesnerie, and Tirole 1985). Second, it captures the realistic scenario that firms can modify concluded contracts only if they do not target specific individuals. For instance, insurance contracts often contain clauses that give the firm the right to modify some terms of the contract, without discriminating between customers and also before a damage event has occurred, leaving the decision whether to accept or to opt out of the contract to the insurant. On the other hand, the insured person can often cancel its policy with relatively short notice.²³ In an education application, long-run contracts that arrange the terms of employment before educational choices have been made often do not exist at all, yielding an equivalent game theoretic structure.

Besides these justifications, there are two additional, methodological reasons for our approach with two-sided and hence a rather strong form of lack of commitment. First, our following results that competitive markets are under some conditions able to deal with the

²³Of course, our model does not predict that such type of behavior must actually be observed in reality. The equilibrium contracts that we derive are renegotiation-proof in the sense that, if they are already offered initially, there will be no incentive for firms to alter them later or for agents to switch to a new contract.

commitment problem rather efficiently becomes stronger the greater the assumed lack of exogenously given commitment opportunities. Second, the assumed structure parallels the one introduced for the social planner above, who is not bound to an initial policy announcement.

Since initial contract offers – if they exist – are not relevant under these assumptions, we can again examine a reduced timing:

Stage 1: Agents choose an effort level.

Stage 2: Some market game takes place, resulting in a set of offered contracts.

Stage 3: Agents simultaneously choose a contract.

The key here is that the comparison between markets and governments that we derive in the following does neither require a detailed specification of a competitive market game in stage 2, nor does it rest on a particular equilibrium notion for the ex post market. We only assume that, after agents have chosen their effort according to a threshold $\hat{d}$, some ex post market game takes place which results in an equilibrium set of contract offers. While the market game for given effort $\hat{d}$ could be complicated, with respect to the timing of moves or its observability assumptions, we restrict attention to its outcome, i.e. to the two contracts among the final offers that maximize the utility of the two different effort types and will thus be chosen in stage 3. We denote this outcome by $V^M(\hat{d}) = (u_{b,h}^M(\hat{d}), u_{b,l}^M(\hat{d}), u_{g,h}^M(\hat{d}), u_{g,l}^M(\hat{d})) \in \mathcal{V}$ and proceed to formulate plausible conditions on $V^M(\hat{d})$, which essentially require informational and resource feasibility as well as a minimal degree of competitive pressure. This approach builds on the insights of Rothschild (2007), who has observed a similar robustness property in a setting of categorical discrimination in insurance markets. In particular, we impose the following axioms:

(C1) $V^M(\hat{d})$ is **incentive compatible**, i.e.

$$p_k u_{k,h}^M(\hat{d}) + (1 - p_k) u_{k,l}^M(\hat{d}) \geq p_k u_{k',h}^M(\hat{d}) + (1 - p_k) u_{k',l}^M(\hat{d}) \quad \forall k, k' \in \{g, b\},$$

(C2) $V^M(\hat{d})$ is **resource feasible**, i.e.

$$G(\hat{d})[p_g \Phi(u_{g,h}^M(\hat{d})) + (1 - p_g) \Phi(u_{g,l}^M(\hat{d}))] + (1 - G(\hat{d}))[p_b \Phi(u_{b,h}^M(\hat{d})) + (1 - p_b) \Phi(u_{b,l}^M(\hat{d}))] \leq \mathbb{E}[\tilde{y}|\hat{d}],$$

(C3) $V^M(\hat{d})$ is **minimally contestable**, i.e. there does not exist an incentive compatible outcome $\tilde{V} = (\tilde{u}_{b,h}, \tilde{u}_{b,l}, \tilde{u}_{g,h}, \tilde{u}_{g,l}) \in \mathcal{V}$ such that

1. $\pi_k(\tilde{u}_{k,h}, \tilde{u}_{k,l}) \geq 0 \quad \forall k \in \{g, b\}$, and

2. $\pi_k(\tilde{u}_{k,h}, \tilde{u}_{k,l}) > 0$ and $p_k \tilde{u}_{k,h} + (1-p_k) \tilde{u}_{k,l} > p_k u_{k,h}^M(\hat{d}) + (1-p_k) u_{k,l}^M(\hat{d})$ for some $k \in \{g, b\}$,

where $\pi_k(u_h, u_l) = p_k(y_h - \Phi(u_h)) + (1-p_k)(y_l - \Phi(u_l))$ are the profits earned with one unit of k -types in contract (u_h, u_l) .

Clearly, any market outcome $V^M(\hat{d})$, whether competitive, monopolistic, or in between, has to satisfy (C1) and (C2). The third requirement (C3), introduced by Rothschild (2007), captures a minimal notion of competition. It implies that a market outcome fails to be minimally contestable only if a firm could offer a pair of incentive compatible deviation contracts that are such that they make non-negative profits no matter what types they attract, and they earn strictly positive profits on some type that strictly prefers them to $V^M(\hat{d})$. This rules out market outcomes that do not survive even the slightest degree of competitive pressure. Let us provide an example for an outcome $V^M(\hat{d})$ that satisfies conditions (C1) to (C3).

Example. Consider the Rothschild-Stiglitz contracts $(u_{g,h}^{RS}, u_{g,l}^{RS})$ and (u_b^{RS}, u_b^{RS}) , where $u_b^{RS} = U(p_b y_h + (1-p_b) y_l)$ is the output-independent payoff for bad types, and the good type's contract satisfies $\pi_g(u_{g,h}^{RS}, u_{g,l}^{RS}) = 0$ and $p_b u_{g,h}^{RS} + (1-p_b) u_{g,l}^{RS} = u_b^{RS}$. These contracts are independent of the threshold $\hat{d}$, and the outcome $V^{RS} = (u_b^{RS}, u_b^{RS}, u_{g,h}^{RS}, u_{g,l}^{RS})$ satisfies conditions (C1) and (C2) for any $\hat{d} \in (0, \infty)$ by definition. It also satisfies (C3) because V^{RS} simultaneously maximizes the utility of both types among the incentive compatible pairs of contracts that break even individually.

The example illustrates that axioms (C1) - (C3) are actually weak. Even though the Rothschild-Stiglitz contracts might not be considered a reasonable market outcome for all $\hat{d} \in (0, \infty)$ (due to equilibrium non-existence problems), they still satisfy the axioms. Also, we will later illustrate that the axioms do not rule out cross-subsidization between ex post types, and outcomes satisfying them might still be susceptible to cream-skimming behavior. But these considerations strengthen our following result, which is based on (C1) - (C3) only, and will thus hold a fortiori for market outcomes that satisfy even stricter requirements.²⁴

²⁴The axioms (C1) - (C3) coincide with those used by Rothschild (2007) except for two differences. First, in the definition of minimal contestability, Rothschild (2007) requires the deviation $\tilde{V}$ to be resource feasible, but this is implied by property 1. in the definition of (C3) and can be omitted. Second, in addition to (C1) - (C3) Rothschild (2007) also requires a market outcome to be *individually rational*, which amounts to the assumption that $p_k u_{k,h}^M(\hat{d}) + (1-p_k) u_{k,l}^M(\hat{d}) \geq p_k U(y_h) + (1-p_k) U(y_l) \forall k \in \{g, b\}$. As we show in the proof of Theorem 1, this axiom is actually not independent, i.e. it is implied by (C1) - (C3) and can also be omitted.

3.3 A Comparison

We will now compare markets based on properties (C1) - (C3) to a planner as analyzed above. We use the following concepts of Pareto dominance. An outcome $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l})$ *ex post Pareto dominates* an outcome $V' = (u'_{b,h}, u'_{b,l}, u'_{g,h}, u'_{g,l})$ if $p_k u_{k,h} + (1 - p_k) u_{k,l} \geq p_k u'_{k,h} + (1 - p_k) u'_{k,l} \forall k \in \{g, b\}$ with strict inequality for at least one $k \in \{g, b\}$, i.e. both ex post effort types weakly prefer their optimal contract in V over their optimal contract in V' , and at least one of them strictly. Because the concept of ex post dominance ignores effort choice, we say that V *ex ante Pareto dominates* V' if for each ex ante cost type $d \in [0, \infty)$, the overall expected utility, including effort cost, under outcome V is weakly larger (and strictly so for at least one d) than the overall utility under V' , when effort is chosen *optimally* for V as well as for V' .²⁵ Based on this terminology, we have the following result:

Theorem 1. *Suppose a market outcome $V^M(\hat{d})$, $\hat{d} \in (0, \infty)$, satisfies (C1)-(C3). Then it ex post Pareto dominates $V^{SP}(0)$. If $V^M(\hat{d})$ also satisfies*

$$\hat{d} = p_g u_{g,h}^M(\hat{d}) + (1 - p_g) u_{g,l}^M(\hat{d}) - p_b u_{b,h}^M(\hat{d}) - (1 - p_b) u_{b,l}^M(\hat{d}), \quad (6)$$

then it also ex ante Pareto dominates $V^{SP}(0)$.

Proof. Suppose $V^M(\hat{d})$ satisfies (C1) - (C3). We first show that $p_k u_{k,h}^M(\hat{d}) + (1 - p_k) u_{k,l}^M(\hat{d}) \geq p_k U(y_h) + (1 - p_k) U(y_l) \forall k \in \{g, b\}$ must hold. Assume to the contrary that this is violated for a type $j \in \{g, b\}$ and consider the outcome $\tilde{V} = (U(y_h) - \epsilon, U(y_l) - \epsilon, U(y_h) - \epsilon, U(y_l) - \epsilon)$ for small $\epsilon > 0$. $\tilde{V}$ is incentive compatible and satisfies $\pi_k(U(y_h) - \epsilon, U(y_l) - \epsilon) > 0 \forall k \in \{g, b\}$ by definition. Also, for ϵ sufficiently small, we have that $p_j(U(y_h) - \epsilon) + (1 - p_j)(U(y_l) - \epsilon) > p_j u_{j,h}^M(\hat{d}) + (1 - p_j) u_{j,l}^M(\hat{d})$ still holds, so that $V^M(\hat{d})$ violates (C3), a contradiction. Thus any outcome that satisfies (C1) - (C3) must be *individually rational* as defined by Rothschild (2007).

Consider the Rothschild-Stiglitz contracts as introduced before. They satisfy $u_{g,h}^{RS} > u_{g,l}^{RS}$ and hence, since $p_g > p_b$, $p_g u_{g,h}^{RS} + (1 - p_g) u_{g,l}^{RS} > u_b^{RS}$. Lemma 4 in Rothschild (2007), considering the special case of only two types, now implies that, for any given $\hat{d} \in (0, \infty)$, the outcome $V^M(\hat{d})$ satisfies $p_g u_{g,h}^M(\hat{d}) + (1 - p_g) u_{g,l}^M(\hat{d}) \geq p_g u_{g,h}^{RS} + (1 - p_g) u_{g,l}^{RS}$ and $p_b u_{b,h}^M(\hat{d}) + (1 - p_b) u_{b,l}^M(\hat{d}) \geq u_b^{RS}$, i.e. both types are ex post weakly better off in $V^M(\hat{d})$ than in the Rothschild-Stiglitz contracts. Since $u_b^{RS} = u^{SP}(0)$ and $p_g u_{g,h}^{RS} + (1 - p_g) u_{g,l}^{RS} > u_b^{RS} = u^{SP}(0)$, this implies that $V^M(\hat{d})$ ex post Pareto dominates $V^{SP}(0)$.

If $\hat{d}$ and $V^M(\hat{d})$ satisfy condition (6), then $p_g u_{g,h}^M(\hat{d}) + (1 - p_g) u_{g,l}^M(\hat{d}) - d > p_b u_{b,h}^M(\hat{d}) + (1 - p_b) u_{b,l}^M(\hat{d}) \geq u_b^{RS} = u^{SP}(0)$ for all $d < \hat{d}$, where the first inequality follows direct from condition (6), and the other comparisons from the above argument for ex post Pareto dominance. Hence under

²⁵The definition is based on optimal effort choice for both outcomes because in any equilibrium, in markets and under a social planner, the agents will actually choose effort optimally in anticipation of the continuation equilibrium contracts.

(6), the threshold $\hat{d}$ describes optimal effort choice for outcome $V^M(\hat{d})$, and all agents who prefer the high effort ($d < \hat{d}$) and subsequently contract $(u_{g,h}^M(\hat{d}), u_{g,l}^M(\hat{d}))$ have a strictly larger utility, including effort cost, than they obtain as bad types in $V^{SP}(0)$. Agents preferring the low effort ($d \geq \hat{d}$) are weakly better off in $V^M(\hat{d})$ compared to $V^{SP}(0)$ from the above ex post Pareto result. $\square$

To be able to make Pareto comparisons, we thus do not need to know exactly how markets work. The theorem shows that whenever the outcome of a market with a fixed interior share of good types satisfies the plausible conditions (C1) to (C3), the contracts that it provides to both types are Pareto better than those implemented by a social planner in the ESP. This ex post comparison is, however, still incomplete for two reasons. First, it ignores the effort costs paid by good types. Second, it compares the planner's performance in the game without commitment to the market's performance with given effort. But suppose that there exists a $\hat{d}^M \in (0, \infty)$ such that the associated outcome $V^M(\hat{d}^M)$ satisfies the fixed point condition (6). Then we can consider $V^M(\hat{d}^M)$ an equilibrium outcome in the market game without commitment, and Theorem 1 shows that the market outcome under lack of commitment Pareto dominates the planner's outcome ex ante, taking effort cost into account.²⁶ To show that there actually exist market outcomes that satisfy conditions (C1) - (C3) and constitute an interior fixed point of (6), we again consider the example from the previous section, which also illustrates that the market outcomes $V^M(\hat{d})$ must involve a sufficient degree of separation for all $\hat{d} \in (0, \infty)$ to guarantee the existence of such a fixed point.

Example continued. Consider the Rothschild-Stiglitz outcome V^{RS} . Since V^{RS} is separating with $p_g u_{g,h}^{RS} + (1 - p_g) u_{g,l}^{RS} > u_b^{RS}$, the optimal critical value for effort choice $\Delta \equiv p_g u_{g,h}^{RS} + (1 - p_g) u_{g,l}^{RS} - u_b^{RS}$ is strictly positive and independent of the exogenously given $\hat{d}$. Thus $\hat{d} = \Delta$ is the (unique) interior value of $\hat{d}$ that satisfies the fixed point condition (6).

More general sufficient conditions for the existence of such a fixed point would, for example, be continuity of $V^M(\hat{d})$ in $\hat{d}$ together with boundedness of $D^M(\hat{d}) \equiv p_g u_{g,h}^M(\hat{d}) + (1 - p_g) u_{g,l}^M(\hat{d}) - p_b u_{b,h}^M(\hat{d}) - (1 - p_b) u_{b,l}^M(\hat{d})$ and $\lim_{\hat{d} \rightarrow 0} D^M(\hat{d}) > 0$, which requires that some separation is sustained as the share of good types goes to zero. Sufficient ex post separation is indeed necessary to obtain a no commitment equilibrium in which good types do exist. But separation and existence of both types does not already guarantee that the corresponding allocation Pareto dominates the planner's outcome. If the minimal notion of competition as captured by condition (C3) is not satisfied by a separating outcome, we cannot expect it to be Pareto superior to $V^{SP}(0)$. Consider, for example, a monopolistic firm that screens the

²⁶See section 4 for a rigorous game-theoretic analysis of equilibria with competitive markets, which also addresses the issue of equilibrium existence.

population of agents. For any given threshold $\hat{d}$ and share of good type $G(\hat{d})$, the monopolist will extract from the agents as many resources as possible.²⁷ Then, even if there is an equilibrium with no commitment in which both types do exist and are separated, there is no reason to expect that its outcome leaves both types better off than in $V^{SP}(0)$.²⁸ Hence, while our comparison between markets and governments does not depend on the details of the market game, it depends on the assumption that the market satisfies a minimal notion of competition.

4 A Game Theoretic Foundation

The results in the previous section were based on plausible assumptions about market outcomes. It remains to be shown that there is indeed a reasonable game in which the discussed properties arise in equilibrium. There are several requirements such a game should satisfy. First, an ex post market equilibrium should exist for every stage 1 effort choice as described by $\hat{d}$. Second, firms should not be restricted to offer only one contract, as such a restriction is not imposed on a social planner either.²⁹ Neither of these conditions are satisfied by the basic model of Rothschild and Stiglitz (1976) that we used to illustrate Theorem 1 above. In this section, we therefore construct an extensive form game in which firms can offer any finite number of contracts, and we characterize sequential equilibria of it.³⁰

4.1 The Market Game

Let $\mathcal{J} = \{0, 1, 2, \dots, N\}$ be a set of risk-neutral firms, each of which can offer up to $r \geq 2$ contracts from $\mathcal{I}$.³¹ We assume that firm 0 is not a regular player of the game but always offers a contract $(\bar{u}_h, \bar{u}_l)$, of which we only assume that it imposes no binding constraint on the optimization problem in Section 4.2. Choosing firm 0's contract corresponds to choosing

²⁷Clearly, the outcome under a monopolist will depend on the specification of the agents' outside options, which we could ignore so far. Since a planner maximizes welfare ex post, she implements contracts on which a reasonable outside option, such as for example the possibility to remain uninsured $(U(y_h), U(y_l))$, imposes no binding constraint. The same conclusion holds for weakly competitive markets from our previous results.

²⁸In fact, the agents will be strictly worse off whenever their outside option is sufficiently unattractive.

²⁹This requirement is of greater importance than it might appear at first glance. Restricting the number of contracts that firms can offer amounts to a restriction on their ex post deviation possibilities. Such a restriction reduces the scope for profitable deviations from initial announcements and thus acts to increase commitment.

³⁰This is in contrast to most approaches in the literature on competitive insurance markets, such as Rothschild and Stiglitz (1976), Wilson (1977) and Hellwig (1987), who restrict a given firm to offer a single contract only. Our game-theoretic approach is also different from the general equilibrium approaches in Dubey and Geanakoplos (2002) and Bisin and Gottardi (2006), although the equilibria we construct are always constrained-efficient as in the latter model.

³¹The equilibria that we construct require the existence of at least 4 active firms.

an outside option, such as carrying out a project without a firm, or remaining uninsured.³² The extensive form is the following:

Stage 1: Agents choose an effort level.

Stage 2a: Firms simultaneously decide on their contract offers.

Stage 2b: After observing all contract offers from stage 2a, firms simultaneously decide whether to remain in the market or to become inactive. Becoming inactive requires all offered contracts to be withdrawn, with a resulting payoff of $-\delta \leq 0$.

Stage 3: Agents simultaneously choose among all remaining offers.

Stages 1 and 3 are as before. Since we do not restrict the number of contracts that a given firm can offer in stage 2a to be one, the presence of another stage 2b in which firms can become inactive is crucial to avoid problems of equilibrium nonexistence that arise otherwise, as already observed by Miyazaki (1977).³³ Unfortunately, the withdrawal phase also introduces a problem of equilibrium multiplicity: Non-competitive equilibria emerge where several firms offer competitive contracts only to withdraw them in equilibrium, but credibly threaten to remain active if they observe deviations in stage 2a. These equilibria are, however, not robust in the sense that they are destroyed by arbitrarily small withdrawal costs. On the other hand, large withdrawal costs would effectively eliminate stage 2b and thus lead to equilibrium nonexistence. These arguments motivate our introduction of withdrawal costs to select *robust* equilibria, i.e. equilibria that exist if withdrawal is costless ($\delta = 0$) but still for sufficiently small values of $\delta > 0$.

We solve the game backwards and begin after stage 1, when effort choices of all agents are given by a critical value $\hat{d}$ as before. From the perspective of the firms, agents then differ only with respect to their effort type, and the game reduces to a standard market with exogenous, private types from stage 2 on. As outlined in Section 2, firms cannot observe effort choice. For the analysis of this and the next subsection, however, let us assume that they know the threshold $\hat{d}$ for effort choice and hence the share $G(\hat{d})$ of good types in the population at the beginning of stage 2. When characterizing equilibria of the entire game, this assumption will

³²Hence a natural outside option, which satisfies our requirement, could be the contract $(\bar{u}_h, \bar{u}_l)$ with $\bar{u}_h = U(y_h)$ and $\bar{u}_l = U(y_l)$. The only role firm 0 plays in the following is to make sure that there is always a non-empty set of contracts the agents can choose from.

³³Our extensive form is different from the suggestion by Miyazaki (1977) and also from the contestable monopoly model by Fernandez and Rasmussen (1993), where firms offer multiple contracts but can withdraw individual contracts. While all of our results remain unchanged if withdrawing individual contracts is possible, the present approach has the advantage that it allows for equilibria with more than just one active firm in the market. Also, while some of our proofs are similar, Fernandez and Rasmussen (1993) resort to a special equilibrium concept, which is not used here.

be motivated by a sequential consistency requirement on the firms' beliefs in equilibrium.

We start with characterizing the agents' optimal strategies in stage 3 for any history of play up to stage 2. First, we restrict the history dependence of agents' strategies such that they are contingent only on the set of offered contracts available after stage 2, excluding the possibility that choices depend on the history of offers and withdrawals. We can then describe contract choices in stage 3 by functions $I_k : \mathcal{Q} \rightarrow \mathcal{I}$, $k = g, b$, that give the contract $I_k(Q) \in Q$ that an agent of ex-post type $k = g, b$ chooses out of any offered set Q . In particular, optimality of choice requires that for $k = g, b$ and any $Q \in \mathcal{Q}$,

$$I_k^*(Q) \in \arg \max_{(u_h, u_l) \in Q} p_k u_h + (1 - p_k) u_l. \quad (7)$$

We restrict attention to stage 3 strategies according to which the contract with smaller difference $u_h - u_l$, that is, with weaker incentives, is chosen in case of indifference. Moreover, we assume that, whenever the optimal contract for an ex post type is offered by several different firms, then each firm receives the same share of these individuals.

We now proceed backwards to stages 2a and 2b by subsuming the agents' optimal strategies in stage 3 into the firms' payoff functions. Stages 2a and 2b then constitute a well-defined extensive form game of complete information, denoted by $\Gamma^{\hat{d}}$, in which firms are the only strategic players (we suppress the dependency of the game on the withdrawal cost parameter δ for notational convenience). Pure strategy profiles are denoted by $s = (s_0, s_1, \dots, s_N) \in S$, where a firm's strategy $s_j = (s_j^1, s_j^2)$ has two components. First, s_j^1 is a set (possibly empty) of up to r contracts to be offered at stage 2a. Let S_j^1 be the set of possible first period offers of firm j .³⁴ Then, $S^1 = \prod_{j \in \mathcal{J}} S_j^1$ is the set of possible histories to be observed at the beginning of stage 2b, and we can associate a stage 2b subgame $\Gamma^{\hat{d}}(\tilde{s}^1)$ to each history $\tilde{s}^1 = (\tilde{s}_0^1, \dots, \tilde{s}_N^1) \in S^1$. For each subgame, s_j^2 then prescribes a withdrawal decision, i.e. $s_j^2 : S^1 \rightarrow \{NW, W\}$, where NW stands for no withdrawal and W for withdrawal. As before, denote by S_j^2 the set of firm j 's possible stage 2 strategies and by $S^2 = \prod_{j \in \mathcal{J}} S_j^2$ the set of stage 2 strategy profiles.³⁵ Then, given a profile $s^2 \in S^2$ of functions, $s^2(\tilde{s}^1) = (s_0^2(\tilde{s}^1), \dots, s_N^2(\tilde{s}^1)) \in \{NW, W\}^{N+1}$ is the vector of withdrawal decisions that s^2 prescribes after history $\tilde{s}^1$.

We first define payoffs for each stage 2b subgame $\Gamma^{\hat{d}}(\tilde{s}^1)$. Let $Q(s^2(\tilde{s}^1)|\tilde{s}^1)$ be the nonempty (due to existence of company 0), finite set of contracts that is available for choice

³⁴Formally, $S_j^1 = \{Q \in \mathcal{Q} | |Q| \leq r\} \cup \emptyset$ for $j = 1, \dots, N$, and $S_0^1 = \{(\bar{u}_h, \bar{u}_l)\}$.

³⁵Formally, $S_j^2 = \{NW, W\}^{S^1}$ for $j = 1, \dots, N$, while S_0^2 is the singleton set containing only the function $s_0^2(\tilde{s}^1) = NW, \forall \tilde{s}^1 \in S^1$.

at the end of $\Gamma^{\hat{d}}(\tilde{s}^1)$ under the withdrawal decisions $s^2(\tilde{s}^1)$. Formally,

$$Q(s^2(\tilde{s}^1)|\tilde{s}^1) = \bigcup_{\substack{j \in \mathcal{J} / \\ s_j^2(\tilde{s}^1) = NW}} \tilde{s}_j^1.$$

As before, let $\pi_k(u_h, u_l)$ denote the profits earned with one unit of k -types in contract (u_h, u_l) . Then, the payoffs of firm j in subgame $\Gamma^{\hat{d}}(\tilde{s}^1)$ are given by $\Pi_j^{\hat{d}}(s^2(\tilde{s}^1)|\tilde{s}^1) = -\delta$ if $s_j^2(\tilde{s}^1) = W$ and otherwise, if $s_j^2(\tilde{s}^1) = NW$, by

$$\Pi_j^{\hat{d}}(s^2(\tilde{s}^1)|\tilde{s}^1) = f\pi_g(I_g^*(Q)) \frac{\mathbf{1}_{\tilde{s}_j^1}(I_g^*(Q))}{\sum_{\substack{i \in \mathcal{J} / \\ s_i^2(\tilde{s}^1) = NW}} \mathbf{1}_{\tilde{s}_i^1}(I_g^*(Q))} + (1-f)\pi_b(I_b^*(Q)) \frac{\mathbf{1}_{\tilde{s}_j^1}(I_b^*(Q))}{\sum_{\substack{i \in \mathcal{J} / \\ s_i^2(\tilde{s}^1) = NW}} \mathbf{1}_{\tilde{s}_i^1}(I_b^*(Q))},$$

where $f = G(\hat{d})$, $Q = Q(s^2(\tilde{s}^1)|\tilde{s}^1)$, and $\mathbf{1}_X$ is the indicator function of set X . Given a strategy profile $s \in S$, the actual payoff of firm j in $\Gamma^{\hat{d}}$ is then $\Pi_j^{\hat{d}}(s) = \Pi_j^{\hat{d}}(s^2(s^1)|s^1)$. Mixed strategies and the associated payoffs are defined analogously.

We are interested in pure strategy subgame perfect equilibria (SPE) of $\Gamma^{\hat{d}}$. However, we have to allow for randomization in some off-equilibrium path stage 2b subgames which do not have a Nash equilibrium in pure strategies.³⁶ We denote equilibrium candidates by $\sigma(\hat{d})$, i.e. strategy profiles that are pure everywhere except in such off-equilibrium path subgames, and SPE of $\Gamma^{\hat{d}}$ by $\sigma^*(\hat{d})$. In stage 3 of any SPE, agents are faced with a nonempty set Q of available contract offers and will make their choices $I_g^*(Q)$ or $I_b^*(Q)$, respectively. Thus, any SPE $\sigma^*(\hat{d})$ is associated with an *outcome* $V^*(\hat{d}) = (u_{b,h}^*(\hat{d}), u_{b,l}^*(\hat{d}), u_{g,h}^*(\hat{d}), u_{g,l}^*(\hat{d})) \in \mathcal{V}$ that summarizes the optimal contract choices of both ex post types in equilibrium.

4.2 A Characterization of SPE Outcomes with Given Effort

In this subsection, we characterize the set of SPE outcomes $V^*(\hat{d})$ of $\Gamma^{\hat{d}}$ for varying levels of δ and any $\hat{d} \in (0, \infty)$. As we will demonstrate below, this set is fundamentally related to the set of solutions to the following optimization problem, which we call problem $\text{GE}(\hat{d})$ (given effort):

$$\max_{(u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}} p_g u_{g,h} + (1-p_g)u_{g,l} \quad (8)$$

³⁶In the SPE that we construct in the proofs of Propositions 2 and 5, any potentially profitable deviation is destroyed using a pure strategy Nash equilibrium in stage 2b. Thus we allow for randomization only to guarantee that there exists a Nash equilibrium in all of the (uncountably many) stage 2b subgames, including those that cannot even be reached by unilateral deviations, but we do not use randomization explicitly in our constructions.

subject to the constraints

$$p_g u_{g,h} + (1 - p_g) u_{g,l} \geq p_g u_{b,h} + (1 - p_g) u_{b,l}, \quad (9)$$

$$p_b u_{b,h} + (1 - p_b) u_{b,l} \geq p_b u_{g,h} + (1 - p_b) u_{g,l}, \quad (10)$$

$$G(\hat{d}) [p_g \Phi(u_{g,h}) + (1 - p_g) \Phi(u_{g,l})] + (1 - G(\hat{d})) [p_b \Phi(u_{b,h}) + (1 - p_b) \Phi(u_{b,l})] \leq \mathbb{E}[\tilde{y}|\hat{d}], \quad (11)$$

$$\Phi(p_b u_{b,h} + (1 - p_b) u_{b,l}) \geq p_b y_h + (1 - p_b) y_l. \quad (12)$$

In program $GE(\hat{d})$, the expected utility of good types is maximized under the ex post incentive compatibility constraints (9) and (10), the resource constraint (11), and a constraint (12) that requires the certainty equivalent of the bad types' contract to be at least as large as their expected endowment. This constraint makes sure that cross-subsidization can only go from good to bad types in any solution, because it implies that the resource cost of the bad types' contract must always be weakly larger than their expected output, i.e. it earns zero or negative profits taken on its own. Note that, comparing with $SP(\hat{d})$, the only two differences are the additional constraint (12) and the objective function (8), which is a special case of (2) putting weight exclusively in good types. The following lemma characterizes the solution to program $GE(\hat{d})$:

Lemma 2. *For any given $\hat{d} \in (0, \infty)$, $GE(\hat{d})$ has a unique solution $V^{GE}(\hat{d}) = (u_{b,h}^{GE}(\hat{d}), u_{b,l}^{GE}(\hat{d}), u_{g,h}^{GE}(\hat{d}), u_{g,l}^{GE}(\hat{d}))$. It is such that $u_{b,h}^{GE}(\hat{d}) = u_{b,l}^{GE}(\hat{d}) \equiv u_b^{GE}(\hat{d})$, the constraints (10) and (11) are binding, and (9) is slack. Moreover, it satisfies conditions (C1) to (C3) from Section 3.2.*

Proof. See Appendix A.2. □

The lemma states that bad types obtain a flat contract in which their utility is output-independent, while the good types' utility depends on their stochastic output, i.e. they are only partially insured. The bad types' incentive compatibility constraint is binding, and resources are exhausted. Moreover, the solution is unique for any given $\hat{d} \in (0, \infty)$.

Let us denote the set of SPE outcomes of $\Gamma^{\hat{d}}$ for given $\hat{d} \in (0, \infty)$ and cost parameter $\delta \geq 0$ by $\Omega^*(\delta, \hat{d}) \subseteq \mathcal{V}$. Then the main result of this subsection is the following:

Proposition 2. *(i) For any $\hat{d} \in (0, \infty)$ and $\delta > 0$, $\Omega^*(\delta, \hat{d}) \subseteq \{V^{GE}(\hat{d})\} \subseteq \Omega^*(0, \hat{d})$.
(ii) Given any $\hat{d} \in (0, \infty)$, there exists a $\bar{\delta} > 0$ such that $V^{GE}(\hat{d}) \in \Omega^*(\delta, \hat{d})$ for all $\delta < \bar{\delta}$.*

Proof. See Appendix A.3. □

Consider a fixed value of $\hat{d} \in (0, \infty)$. Part (i) of the proposition states that, whenever withdrawal costs are strictly positive, the set of SPE outcomes is either empty due to equilibrium nonexistence (which will be the case if withdrawal costs are too high), or it contains exactly the solution to $\text{GE}(\hat{d})$.³⁷ This solution is also an equilibrium outcome if $\delta = 0$, but additional equilibria with new outcomes can emerge in this case. All these additional outcomes are, however, not robust, i.e. they disappear for arbitrarily small withdrawal costs $\delta > 0$. On the other hand, part (ii) states that the solution to $\text{GE}(\hat{d})$ is actually robust, because it is an SPE outcome for sufficiently small but strictly positive values of δ . The proposition implies that, for any sequence $\{\delta_n\}_{n=0}^\infty$ with $\delta_n \geq 0 \forall n$ and $\lim_{n \rightarrow \infty} \delta_n = 0$,

$$\limsup_{n \rightarrow \infty} \Omega^*(\delta_n, \hat{d}) = \liminf_{n \rightarrow \infty} \Omega^*(\delta_n, \hat{d}) = \{V^{GE}(\hat{d})\},$$

where

$$\limsup_{n \rightarrow \infty} \Omega^*(\delta_n, \hat{d}) = \bigcap_{n=1}^\infty \bigcup_{m=n}^\infty \Omega^*(\delta_m, \hat{d}) \quad \text{and} \quad \liminf_{n \rightarrow \infty} \Omega^*(\delta_n, \hat{d}) = \bigcup_{n=1}^\infty \bigcap_{m=n}^\infty \Omega^*(\delta_m, \hat{d}).$$

In an insurance setting, Proposition 2 can be interpreted as showing that our market game $\Gamma^{\hat{d}}$ produces Miyazaki-Wilson type contracts as the unique robust equilibrium outcome. There are two cases depending on whether constraint (12) does or does not bind in an SPE outcome. If it does, each contract individually makes zero profits, and we obtain the classical Rothschild-Stiglitz outcome. Otherwise, the full insurance contract makes negative and the partial insurance contract makes positive profits, so that there is cross-subsidization from low to high risks. We will derive comparative static effects of a changing value of $\hat{d}$ on the outcome $V^{GE}(\hat{d})$ in the following subsection when analyzing sequential equilibria of the entire game without commitment. It will turn out there that cross-subsidization occurs whenever $\hat{d}$ and thus the share of good types $G(\hat{d})$ is large, while the Rothschild-Stiglitz contracts are the equilibrium outcome for small values of $\hat{d}$.

For completeness, we briefly turn to the case where $\hat{d} \in \{0, \infty\}$, so that all agents are of the same ex post type and there is no issue of asymmetric information. First, define $V^{GE}(\hat{d}) = (u_{b,h}^{GE}(\hat{d}), u_{b,l}^{GE}(\hat{d}), u_{g,h}^{GE}(\hat{d}), u_{g,l}^{GE}(\hat{d}))$ by $u_{b,h}^{GE}(\hat{d}) = u_{b,l}^{GE}(\hat{d}) = u_{g,h}^{GE}(\hat{d}) = u_{g,l}^{GE}(\hat{d}) = U(\mathbb{E}[\tilde{y}|\hat{d}])$ for $\hat{d} \in \{0, \infty\}$. Next, it is straightforward to show that SPE $\sigma^*(\hat{d})$ of $\Gamma^{\hat{d}}$ do exist for $\hat{d} \in \{0, \infty\}$, irrespective of the value of $\delta \geq 0$. As in Section 3, we still write SPE outcomes as elements of $\mathcal{V}$, denoted by $V^*(\hat{d})$, and it follows quickly that for all $\delta > 0$, $\Omega^*(\delta, \hat{d}) = \{V^{GE}(\hat{d})\} \subseteq \Omega^*(0, \hat{d})$ if $\hat{d} \in \{0, \infty\}$. Put differently, whenever withdrawal costs

³⁷Observe that this is a statement about uniqueness of the equilibrium outcome, not about the equilibrium itself. There are always multiple equilibria with the same outcome, which differ with respect to irrelevant contract offers that no agents chooses.

are strictly positive, in any outcome of any SPE, the unique ex post type obtains a contract with a flat payment that corresponds to its expected output. If $\delta = 0$, additional equilibria emerge even if $\hat{d} \in \{0, \infty\}$, but none of them is robust.³⁸

4.3 Sequential Equilibria Without Commitment

We now proceed to analyzing equilibria of the entire market game without commitment, including optimal effort choice by the agents, defined as follows.

Definition 2. *An equilibrium with no commitment (ENC) is a pair $(\hat{d}, \sigma^*(\hat{d}))$ where*

- (i) $\sigma^*(\hat{d})$ is an SPE of $\Gamma^{\hat{d}}$ and*
- (ii) $\hat{d} = p_g u_{g,h}^*(\hat{d}) + (1 - p_g) u_{g,l}^*(\hat{d}) - p_b u_{b,h}^*(\hat{d}) - (1 - p_b) u_{b,l}^*(\hat{d})$.*

The definition is based on the following reasoning. Assume that agents' effort choices are given by a threshold $\hat{d}$, which is without loss of generality as argued before. Firms then observe neither individual nor aggregate effort but have to form a belief about $\hat{d}$, which has to be correct on the equilibrium path. We are interested in sequential equilibria of the subsequent game of incomplete information between insurers, that is, we impose the sequential consistency condition that requires the beliefs to be correct even after stage 2a deviations, when companies observe contract offers which do not occur in equilibrium. Firms also need to behave sequentially rational, i.e. they must make optimal choices in all information sets. But then, any profile of firms' strategies for the game of incomplete information that simultaneously satisfies sequential consistency and rationality is equivalent to a SPE of $\Gamma^{\hat{d}}$, the game with given effort described in the previous subsections.³⁹ Given that agents in stage 1 anticipate the outcome of this game, their actually optimal effort choice is prescribed by a threshold that equals the right-hand side of (ii), and there is no incentive to deviate if and only if the fixed point condition (ii) in the definition is satisfied.⁴⁰

For any given value of $\delta \geq 0$, denote by $\hat{\Omega}^*(\delta)$ the set of ENC outcomes. We again

³⁸Besides the different definition of $V^{GE}(\hat{d})$, the only difference to the previous result is therefore that large withdrawal costs do not lead to equilibrium nonexistence if there is no asymmetric information.

³⁹Of course, our definition is slightly abusive, as the game with incomplete information has non-singleton information sets while all information sets are singletons in the game with given effort. But since there is a bijective relationship between the firms' information sets in the two games, we use strategy profiles σ of the game $\Gamma^{\hat{d}}$ to represent the corresponding strategies in the game of incomplete information.

⁴⁰Analyzing sequential equilibria of a game without commitment, our approach follows those in Chari and Kehoe (1990) and Phelan and Stacchetti (2001), although in a different framework. Both of these papers consider a time-inconsistency problem related to capital taxation as opposed to moral hazard, and do not consider competitive markets. Chari and Kehoe (1990) refer to their equilibria as 'sustainable equilibria', but they are in fact sequential equilibria as observed by Phelan and Stacchetti (2001).

characterize the robust ENC outcomes. For this purpose, it is useful to define the set

$$\Omega^{NC} \equiv \{V^{GE}(\hat{d}) | \hat{d} = p_g u_{g,h}^{GE}(\hat{d}) + (1 - p_g) u_{g,l}^{GE}(\hat{d}) - p_b u_{b,h}^{GE}(\hat{d}) - (1 - p_b) u_{b,l}^{GE}(\hat{d})\} \subset \mathcal{V}.$$

Ω^{NC} contains those robust SPE outcomes of the game with given effort for which the exogenously given $\hat{d}$ constitutes a fixed point in the sense of Definition 2. The first step to characterizing robust ENC outcomes is the following proposition.

Proposition 3. (i) For any $\delta > 0$, $\hat{\Omega}^*(\delta) \subseteq \Omega^{NC} \subseteq \hat{\Omega}^*(0)$.

(ii) For each $V \in \Omega^{NC}$, there exists a $\bar{\delta} > 0$ such that $V \in \hat{\Omega}^*(\delta)$ for all $\delta < \bar{\delta}$.

Proof. We first show that, if $\delta > 0$, $\hat{\Omega}^*(\delta) \subseteq \Omega^{NC}$. Consider any $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \hat{\Omega}^*(\delta)$ and define $\bar{d} \equiv p_g u_{g,h} + (1 - p_g) u_{g,l} - p_b u_{b,h} - (1 - p_b) u_{b,l}$. By definition of ENC, it must then be true that V is an SPE outcome of $\Gamma^{\bar{d}}$, i.e. $V \in \Omega^*(\delta, \bar{d})$. Proposition 2 (or the analogous arguments if $\bar{d} \in \{0, \infty\}$) then implies that $V = V^{GE}(\bar{d})$, which immediately implies that $V \in \Omega^{NC}$, and hence $\hat{\Omega}^*(\delta) \subseteq \Omega^{NC}$.

Second, for any $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \Omega^{NC}$ there exists a $\bar{d}$ such that $V = V^{GE}(\bar{d})$ and $\bar{d} = p_g u_{g,h} + (1 - p_g) u_{g,l} - p_b u_{b,h} - (1 - p_b) u_{b,l}$. It then follows from Proposition 2 (or the analogous arguments if $\bar{d} \in \{0, \infty\}$) that there exists a value $\bar{\delta}$ such that $V \in \Omega^*(\delta, \bar{d})$ for all $0 \leq \delta < \bar{\delta}$, and thus $V \in \hat{\Omega}^*(\delta)$ for all $0 \leq \delta < \bar{\delta}$. This implies statement (ii) and also that $\Omega^{NC} \subseteq \hat{\Omega}^*(0)$. $\square$

According to the proposition, the set of robust ENC outcomes coincides with Ω^{NC} . While any other potential ENC outcome can be destroyed by the introduction of arbitrarily small withdrawal costs, any element of Ω^{NC} is actually an ENC outcome for sufficiently small but positive δ . This again implies that the set of sequential equilibrium outcomes $\hat{\Omega}^*(\delta)$ converges to Ω^{NC} as δ vanishes. Formally, given a sequence $\{\delta_n\}_{n=0}^{\infty}$ with $\delta_n \geq 0 \forall n$ and $\lim_{n \rightarrow \infty} \delta_n = 0$, $\limsup_{n \rightarrow \infty} \hat{\Omega}^*(\delta_n) = \liminf_{n \rightarrow \infty} \hat{\Omega}^*(\delta_n) = \Omega^{NC}$.

Since the definition of the set Ω^{NC} itself relies on a fixed point condition, it is useful for the subsequent analysis to explicitly define a function D as follows. For any given threshold $\hat{d} \in [0, \infty) \cup \{\infty\}$ for effort choice, consider the unique robust SPE outcome $V^{GE}(\hat{d})$ of $\Gamma^{\hat{d}}$. The contracts given by $V^{GE}(\hat{d})$ then induce the critical value

$$D(\hat{d}) = p_g u_{g,h}^{GE}(\hat{d}) + (1 - p_g) u_{g,l}^{GE}(\hat{d}) - p_b u_{b,h}^{GE}(\hat{d}) - (1 - p_b) u_{b,l}^{GE}(\hat{d})$$

for optimal effort choice by all agents if they anticipate $V^{GE}(\hat{d})$, and Ω^{NC} can be rewritten as $\Omega^{NC} = \{V^{GE}(\hat{d}) | \hat{d} = D(\hat{d})\}$. If $\hat{d} \in (0, \infty)$, we know from Lemma 2 that $u_b^{GE}(\hat{d}) = p_b u_{b,h}^{GE}(\hat{d}) + (1 - p_b) u_{b,l}^{GE}(\hat{d})$, so that we can simplify $D(\hat{d})$ to

$$D(\hat{d}) = (p_g - p_b) \left(u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d}) \right) \quad (13)$$

in this case. If $\hat{d} \in \{0, \infty\}$, i.e. all agents are either good or bad types, we immediately obtain $D(0) = D(\infty) = 0$. Since the continuation contract is an output-independent payment in this case (or, equivalently, a full insurance contract), nobody finds it optimal to invest effort to increase the probability of high output.

Let us collect some useful properties of the function D in the following lemma. These properties are based on comparative static effects of varying levels of $\hat{d}$ on the robust SPE outcome $V^{GE}(\hat{d})$.

Lemma 3. (i) D is continuous in $(0, \infty)$.

(ii) $\lim_{\hat{d} \rightarrow 0} D(\hat{d}) > 0$ and $\lim_{\hat{d} \rightarrow \infty} D(\hat{d}) = 0$.

(iii) If

$$\frac{d}{du} \frac{\Phi''(u)}{\Phi'(u)} \geq 0, \quad (14)$$

then there exists $\tilde{d} \in (0, \infty)$ such that $D(\hat{d})$ is flat in $(0, \tilde{d}]$, and strictly decreasing in $\hat{d}$ for all $\hat{d} > \tilde{d}$.

Proof. See Appendix A.4. □

Properties (i) and (ii) together with the fact that $D(0) = 0$ imply that, while D is continuous otherwise, there exists a discontinuity at $\hat{d} = 0$. This is because a contract with output-independent utilities and hence no incentive for effort provision is the unique outcome if $\hat{d} = 0$, while for any positive $\hat{d}$ and also in the limit as $\hat{d} \rightarrow 0$, the good type's contract remains high-powered. Specifically, we show in the proof of the lemma that (12) is binding in $V^{GE}(\hat{d})$ for sufficiently small but positive $\hat{d}$, which is saying that the Rothschild-Stiglitz contracts obtain if the given share of good types is small. As $\hat{d} \rightarrow \infty$, on the other hand, the good type's contract converges to an output-independent, full insurance contract, which requires cross-subsidization to the bad types to preserve incentive-compatibility. Property (iii) shows that, if preferences satisfy condition (14), there is exactly one critical value $\tilde{d}$ at which the transition from zero to positive cross-subsidization occurs. Furthermore, an increase in $\hat{d}$ above $\tilde{d}$ then leads to an increased subsidy and lower-powered incentives $u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d})$, such that D is strictly decreasing.

The following properties of the utility function $U(c)$ are sufficient to guarantee (14):

Lemma 4. (i) If $U(c)$ has constant or increasing absolute risk aversion, (14) is satisfied.

(ii) If $U(c)$ has constant relative risk aversion with coefficient α , then (14) is satisfied if and only if $\alpha \geq 1$.

Proof. See Fudenberg and Tirole (1990), Lemma 3.2. □

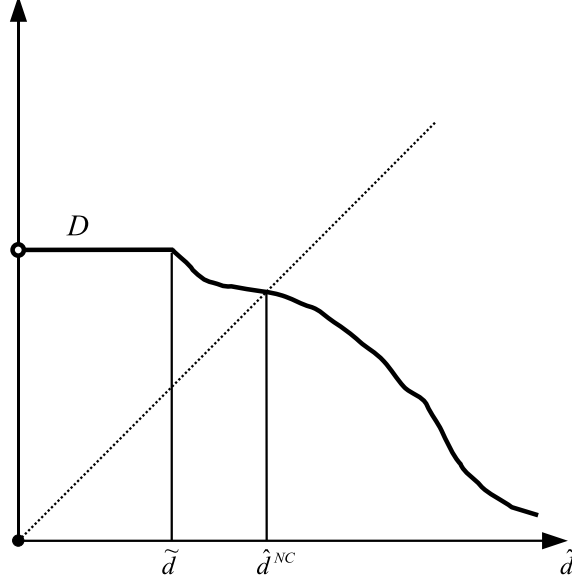


Figure 1: Fixed Point Problem with Competitive Insurance Markets

Figure 1 depicts the function $D(\hat{d})$ for preferences as described in Lemma 4. By definition of D in (13), the shape of $D(\hat{d})$ depends on whether the good types obtain a more or less high-powered incentive contract in response to an increase in $\hat{d}$, as this determines ex ante incentives for effort choice. In the proof of Lemma 3 we show that, given slackness of (12), the robust equilibrium contracts for given $\hat{d}$ solve the following first-order condition of problem $GE(\hat{d})$:

$$\frac{G(\hat{d})}{1 - G(\hat{d})} = \frac{p_g - p_b}{p_g(1 - p_g)} \frac{\Phi'(u_b^{GE}(\hat{d}))}{\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))}. \quad (15)$$

Equation (15) captures the tradeoff between providing more insurance to good types and increasing the cross-subsidy for bad types to preserve incentive compatibility. Condition (14) makes sure that an increase in $\hat{d}$ increases the good types' demand for insurance, $\Phi'(u_{g,h}) - \Phi'(u_{g,l})$, faster than it increases the cost of increasing the bad types' utility, $\Phi'(u_b)$. Then, the good types always obtain a less high-powered contract in response to an increase in $\hat{d}$. For this to be the case, risk aversion must not decline too quickly, which explains the conditions in Lemma 4. Otherwise, the power of the good types' contract may be increasing in $\hat{d}$ in some range, and thus D may have increasing parts.

We can now state the main result of this subsection, which is a direct implication of the previous results and standard fixed point theorems.

Proposition 4. (i) $V^{GE}(0) \in \Omega^{NC}$.

(ii) There exists a value $\hat{d}^{NC} > 0$ such that $V^{GE}(\hat{d}^{NC}) \in \Omega^{NC}$.

(iii) Under condition (14), $\hat{d}^{NC}$ is unique, so that $\Omega^{NC} = \{V^{GE}(0), V^{GE}(\hat{d}^{NC})\}$.

Clearly, $\hat{d} = 0$ is always a fixed point of D , meaning that, for small δ , there is a sequential equilibrium in our model without commitment where no agent exerts effort and everyone obtains a contract with an output-independent payment that corresponds to the level of expected output for bad types. However, there always exists at least one other fixed point $\hat{d}^{NC} > 0$ of D and hence a robust ENC outcome in which a non-zero mass of agents exert the high effort. Since D is weakly decreasing under condition (14), the positive fixed point is unique in this case. Otherwise, multiple non-zero fixed points and associated robust ENC outcomes may exist.

4.4 A Comparison with Governments

We are now in a position to compare welfare in the robust ENC outcomes implemented by competitive markets to that achieved by the government outcomes considered in section 3:

Theorem 2. *Any robust ENC outcome $V^{GE}(\hat{d}^{NC})$ with $\hat{d}^{NC} > 0$ ex-ante Pareto dominates $V^{SP}(0)$, and strongly so if $V^{GE}(\hat{d}^{NC})$ satisfies (12) with slack.*

Proof. We consider two cases, depending on whether constraint (12) is binding or not in $V^{GE}(\hat{d}^{NC})$. Assume first that it does, implying $u_b^{GE}(\hat{d}^{NC}) = U(p_b y_h + (1 - p_b) y_l) = u^{UP}(0)$, i.e. bad types in $V^{GE}(\hat{d}^{NC})$, who do not exert any effort, obtain the same utility as all agents in $V^{UP}(0)$, where nobody exerts any effort. By definition of $\hat{d}^{NC}$, we then have that $p_g u_{g,h}^{GE}(\hat{d}^{NC}) + (1 - p_g) u_{g,l}^{GE}(\hat{d}^{NC}) - d > u_b^{GE}(\hat{d}^{NC}) = u^{UP}(0)$ for all $d < \hat{d}^{NC}$. Given that $\hat{d}^{NC}$ is a fixed point of D , the critical value $\hat{d}^{NC}$ determines optimal effort choice in $V^{GE}(\hat{d}^{NC})$, so that all good types in $V^{GE}(\hat{d}^{NC})$ are ex-ante (including effort cost) strictly better off than they are as bad types in $V^{SP}(0)$. If (12) is slack in $V^{GE}(\hat{d}^{NC})$, then $u_b^{GE}(\hat{d}^{NC}) > u^{SP}(0)$, and both low and bad types are ex-ante strictly better off in $V^{GE}(\hat{d}^{NC})$, with the same argument. $\square$

The first part of Theorem 2 is a corollary of previous results (although we present a convenient direct proof above). From Lemma 2 we know that $V^{GE}(\hat{d}^{NC})$ satisfies (C1) - (C3) if $\hat{d}^{NC} > 0$. The definition of ENC then implies that $V^{GE}(\hat{d}^{NC})$ also satisfies the fixed point condition (6), so that the ex ante Pareto dominance result follows from Theorem 1. We thus know that all agents are weakly and some are strictly better off in competitive markets than under a benevolent government that is concerned about a welfare criterion from a large class that includes the utilitarian case. This holds even from an ex-ante perspective, i.e. taking effort cost into account. Note that Theorem 2 applies to every robust ENC outcome with a positive share of good types, not only if this outcome is unique as under condition (14). The second part of Theorem 2 goes beyond the general insight of Theorem 1. If

the market equilibrium involves cross-subsidization from agents who choose to become good types to those who prefer to become bad types, as captured by slackness of (12), even all agents are strictly better off in the market than under a social planner.

The general intuition for Theorem 2 is that, with a social planner as specified in Proposition 1, all agents end up being bad types in a contract with an output-independent payoff, equal in size to their expected output. In a market, only some agents remain bad types, obtaining a flat payment at the same or even a subsidized level, whereas the other agents prefer to become good types and choose an incentive contract that makes them strictly better off. Thus it is the ex post adverse selection problem in competitive markets, leading to underinsurance of some agents, which is crucial for the dominance of markets over governments.

In an environment with lack of commitment, competitive markets have a clear advantage over a benevolent utilitarian government or one that puts overproportional weight on high effort cost types. Because competitive forces lead to a separating outcome that gives good type agents the best incentive-compatible and resource-feasible contract, better incentives for effort provision can be sustained by markets than by a social planner. At the same time, the cross-subsidization constraint (12) implies that ex post bad types are never worse off in the market than under a centralized regime.

Establishing a competitive market can thus be interpreted as the delegation of allocation decisions to a planner who cares only for high effort types, and who faces the additional cross-subsidization constraint. The similar idea that the creation of an independent agency and the subsequent delegation of decisions to this agency can be beneficial in the presence of a commitment problem is central to the research on central bank independence. In our model, the advantage of a competitive market is that it acts as if it was a specific planner while individual firms are still maximizing profits and hence their real objective. In an independent agency, such as a central bank, there is still a problem of aligning individual incentives with the imposed objective, giving rise to additional moral hazard problems (see, e.g. Walsh (1995)).

5 Extensions

5.1 Ex Post Market Shutdown

Our main results are about the comparison of a centralized solution to a competitive market, without asking how a market comes into being. The concluding discussion of the previous section suggests that one might actually think of a meta-game with an initial phase of

institutional choice, where allocation decisions are delegated to either the government or to a market. In such a game, whenever a market has been established, the question arises whether the government will keep the promise of leaving the market in place ex post. After all, even though the market Pareto dominates the centralized solution from an ex ante perspective, a planner would still be tempted to abolish the market ex post to implement a pooling solution. While it is reasonable to assume that the existence of a market imposes some obstacle on the government's ability to redistribute ex post, we still want to address the question what happens if the government is able to shut down the market with some exogenous probability $1 - q \in [0, 1]$. Will the Pareto-comparison from the previous section generalize to such a situation?

To answer this question, consider again the situation after stage 1, when agents have chosen their effort according to some threshold $\hat{d}$ and the market begins to operate. The assumption that q is exogenous implies that neither effort choice nor the subsequent market outcome can influence the probability that the market will persist.⁴¹ The market will then work exactly as described in the preceding section, producing the outcome $V^{GE}(\hat{d})$. With probability q , this will be the final outcome. Otherwise, assuming $\Psi \succeq_{FOSD} G$ as before, the government steps in and provides full insurance to all agents at the identical level $U(\mathbb{E}[\tilde{y}|\hat{d}])$. The agents anticipate this when making their ex ante effort decisions. It is straightforward to show that the threshold for optimal effort choice in stage 1 is now given by

$$D^q(\hat{d}) = q \left[p_g u_{g,h}^{GE}(\hat{d}) + (1 - p_g) u_{g,l}^{GE}(\hat{d}) - p_b u_{b,h}^{GE}(\hat{d}) - (1 - p_b) u_{b,l}^{GE}(\hat{d}) \right].$$

Obviously, the family of functions $\{D^q\}_{q \in [0,1]}$ contains D as defined in the previous section for the special case when $q = 1$ and markets persist definitely. For $q = 0$, on the other hand, D^q coincides with D^{SP} from Section 3, where the planner always decides on the final allocation. In all intermediate cases $q \in (0, 1)$, D^q is a scaled-down version of D , and Lemma 3 immediately implies that D^q has at least one strictly positive fixed point $\hat{d}^q > 0$. Theorem 2 then still applies, which shows that the market outcome Pareto dominates the centralized solution whenever there is some positive probability that the market is not destroyed by the government, even if this probability is arbitrarily small.⁴²

On the other hand, the possibility of ex post government intervention has an impact on the equilibrium level of effort. In general, agents' ex ante incentives for effort are weakened compared to the case where $q = 1$, since the separating market outcome is realized ex post

⁴¹This is clearly a simplifying assumption. If individual effort choice or the market outcome affect the probability of market shutdown, additional strategic effects appear which are not taken into account in our game theoretic model from the previous section.

⁴²Furthermore, the Pareto comparison extends to expected utilities taking into account the possible shut-down case.

only with some probability, whereas government equalizes consumption otherwise. Assume, for example, that (14) is satisfied. The unique interior fixed point $\hat{d}^q > 0$ of D^q is then increasing in q , such that an increasing probability of market shutdown decreases equilibrium effort.⁴³

5.2 A Social Planner with General Pareto-Weights

In sections 3 and 4, we have obtained strong results comparing welfare under competitive markets and a social planner if the latter is concerned about a utilitarian welfare criterion or aims at redistributing towards agents with high effort cost. If the government is putting overproportional Pareto-weights on low effort cost types compared to population shares, then such clear *welfare* comparisons with markets are not generally available. However, we show in this subsection that it is possible to compare the level of *effort* implemented by a social planner with a general distribution $\Psi(d)$ of Pareto-weights to that implemented by competitive markets.

A first issue that arises in the analysis for general weights is that the function D^{SP} as defined in Section 3 is no longer simple in shape. Depending on Ψ , it can have increasing and decreasing parts, resulting in the possibility of multiple fixed points and thus ESP.⁴⁴ Such equilibria will differ in the aggregate level of effort described by their threshold level $\hat{d}^{SP}$ for effort choice. Despite the potential multiplicity, we still have the following result:

Theorem 3. *Suppose condition (14) is satisfied and $\hat{d}^{NC} > \tilde{d}$. Then $\hat{d}^{SP} \leq \hat{d}^{NC}$ for any ESP $(\hat{d}^{SP}, V^{SP}(\hat{d}^{SP}))$.*

Proof. Fix a value $\hat{d} \in (0, \infty)$ and consider two cases. First, suppose that $\Psi(\hat{d}) \leq G(\hat{d})$. Then, arguing as for the proof of Proposition 1, we obtain $D^{SP}(\hat{d}) = 0$. Second, consider the case $\Psi(\hat{d}) > G(\hat{d})$. Lemma 1 implies that the unique solution to SP($\hat{d}$) is such that constraints (4) and (5) bind and $u_{b,h}^{SP}(\hat{d}) = u_{b,l}^{SP}(\hat{d})$. From the fact that $p_g > p_b$ and $u_{g,h}^{SP}(\hat{d}) \geq u_{g,l}^{SP}(\hat{d})$ is then also follows that (3) is automatically satisfied. Defining

$$p(\hat{d}) \equiv p_g \Psi(\hat{d}) + p_b(1 - \Psi(\hat{d})), \quad (16)$$

the solution must therefore be such that $(u_{g,h}^{SP}(\hat{d}), u_{g,l}^{SP}(\hat{d}))$ solves the simplified problem

$$\max_{(u_{g,h}, u_{g,l}) \in \mathcal{I}} p(\hat{d}) u_{g,h} + (1 - p(\hat{d})) u_{g,l} \quad (17)$$

⁴³If (14) is not satisfied and there are several interior fixed points, the usual comparative statics results apply, i.e. some of the fixed points will be increasing (including the highest one) and others will be decreasing in q .

⁴⁴If Ψ is not continuous, D^{SP} can have discontinuities as well. The existence of an ESP is still guaranteed, because $D^{SP}(0) = 0$ always holds.

subject to the resource constraint

$$G(\hat{d}) [p_g \Phi(u_{g,h}) + (1 - p_g) \Phi(u_{g,l})] + (1 - G(\hat{d})) \Phi(p_b u_{g,h} + (1 - p_b) u_{g,l}) = \mathbb{E}[\tilde{y} | \hat{d}]. \quad (18)$$

Suppose $\hat{d} > \tilde{d}$. First, for $\hat{d} = \infty$, $D^{SP}(\hat{d}) = D(\hat{d}) = 0$ holds. Otherwise, if $\hat{d} \in (\tilde{d}, \infty)$, by Lemma 2 the robust equilibrium outcome with given effort choice in competitive markets $V^{GE}(\hat{d})$ is such that $(u_{g,h}^{GE}(\hat{d}), u_{g,l}^{GE}(\hat{d}))$ solves

$$\max_{(u_{g,h}, u_{g,l}) \in \mathcal{I}} p_g u_{g,h} + (1 - p_g) u_{g,l} \quad (19)$$

subject to the same budget constraint (18), because (12) does not bind for $\hat{d} > \tilde{d}$ as shown in the proof of Lemma 3. Since (18) is a convex constraint by the proof of Lemma 1, and $p(\hat{d}) \leq p_g$ holds by (16), the solutions $(u_{g,h}^{SP}(\hat{d}), u_{g,l}^{SP}(\hat{d}))$ and $(u_{g,h}^{GE}(\hat{d}), u_{g,l}^{GE}(\hat{d}))$ must be such that

$$u_{g,h}^{SP}(\hat{d}) \leq u_{g,h}^{GE}(\hat{d}) \quad \text{and} \quad u_{g,l}^{SP}(\hat{d}) \geq u_{g,l}^{GE}(\hat{d}).$$

This implies

$$D^{SP}(\hat{d}) = (p_g - p_b)(u_{g,h}^{SP}(\hat{d}) - u_{g,l}^{SP}(\hat{d})) \leq (p_g - p_b)(u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d})) = D(\hat{d}). \quad (20)$$

Under property (14) and $\hat{d}^{NC} > \tilde{d}$, the function D is strictly decreasing in $\hat{d}$ above its unique fixed point $\hat{d}^{NC}$, so that $D(\hat{d}) < \hat{d}$ for all $\hat{d} > \hat{d}^{NC}$. Together with (20) and $D^{SP}(\infty) = D(\infty) = 0$, this implies $D^{SP}(\hat{d}) \leq D(\hat{d}) < \hat{d}$ for all $\hat{d} > \hat{d}^{NC}$, so that any fixed point $\hat{d}^{SP}$ of D^{SP} must satisfy $\hat{d}^{SP} \leq \hat{d}^{NC}$. $\square$

According to the theorem, whenever condition (14) is satisfied and the market outcome without commitment involves cross-subsidization from good to bad types (as captured by $\hat{d}^{NC} > \tilde{d}$), the associated equilibrium share of good types $G(\hat{d}^{NC})$ is higher than the share of good types $G(\hat{d}^{SP})$ in any equilibrium with a social planner, irrespective of the distribution of Pareto-weights $\Psi(d)$ that is used. Hence, in terms of incentives for effort, a social planner can do only as well as or worse than competitive markets, but never better. Figure 2 illustrates this comparison. Although D^{SP} may be non-monotonic and have several fixed points, we show in the proof of Theorem 3 that $D^{SP}(\hat{d})$ must always lie below $D(\hat{d})$ for all $\hat{d} \geq \tilde{d}$ if condition (14) is satisfied, which implies the result. The restriction to the case that the market equilibrium involves cross-subsidization is necessary because, while markets are constrained by the fact that there cannot be cross-subsidization from bad to good types in equilibrium, the planner is not. A planner who is otherwise similar to the market, in that she uses a large weight $\Psi(\hat{d}) > G(\hat{d})$ close to one, may find it optimal to provide even stronger incentives for effort provision than the market. This further slackens the incentive constraint and makes it possible to extract additional resources from the bad types, which

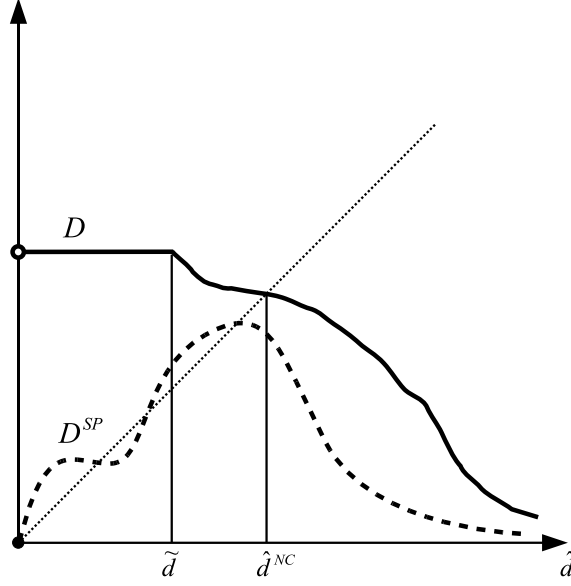


Figure 2: Fixed Point Problem with a Non-Utilitarian Social Planner

could more than compensate the good types for their reduced insurance. As Proposition 1 makes clear, this can only occur if the planner puts sufficiently overproportional welfare weight on low effort cost types. Otherwise, the reason why competitive markets are more successful in providing incentives is again that they replicate an extreme planner who cares only about ex post good types.

5.3 Competitive Markets with Full Commitment

In this subsection, we take another approach to evaluating the performance of competitive markets without commitment. Rather than comparing them to another institution such as a government, we now ask how the equilibrium outcomes studied so far compare to the benchmark of full commitment. We demonstrate that competitive equilibria without and with full commitment fall together whenever the latter involve no cross-subsidization between ex post effort types. This illustrates from yet another perspective that competitive markets are an institution that is able to deal effectively with lack of commitment, even to the degree that the commitment problem may entirely vanish.⁴⁵

⁴⁵Moreover, analyzing markets with full commitment in our model economy is in itself an interesting exercise, as the working of markets in the presence of both adverse selection from private cost parameters and moral hazard from hidden effort choice is not yet well understood, even in a setting without commitment problems. To our knowledge, the only contributions addressing both issues simultaneously, in the framework of insurance markets, are Stewart (1994) and Chassagnon and Chiappori (1997). Their models, however, include additional dimensions of individual heterogeneity, leading to multidimensional screening problems.

Let us consider the following modified timing:

Stage 1a: Firms simultaneously decide on their contract offers.

Stage 1b: After observing all offers, firms simultaneously decide whether to remain active.

Stage 2: Agents choose among remaining offers and between the effort levels.

The only difference to the extensive forms considered in sections 3 and 4 is that effort choice takes place *after* binding contract offers have been announced. Firms are therefore able to take into account the effect of their contract offers on the composition of their customer pool in terms of effort types, rather than taking it as given.

As in Section 4, we again subsume the agents' optimal strategies for effort and contract choice in stage 2 into the firms' payoff functions and consider stages 1a and 1b as a dynamic game of complete information between firms, which we denote by Γ^{FC} . Firms' strategies and strategy profiles are exactly as in Section 4, and payoffs are defined analogously. Then, any SPE σ^* of Γ^{FC} is again associated with an outcome $V^* = (u_{b,h}^*, u_{b,l}^*, u_{g,h}^*, u_{g,l}^*)$ representing the best contracts for both effort types. In addition, the threshold for effort choice in V^* is

$$\hat{d}^* = p_g u_{g,h}^* + (1 - p_g) u_{g,l}^* - p_b u_{b,h}^* - (1 - p_b) u_{b,l}^*.$$

In principle we could have $\hat{d}^* = \infty$ or $\hat{d}^* \leq 0$, i.e. all agents might decide to become of the same ex post type.⁴⁶ Observe that the outcome V^* is then still well-defined, because an optimal contract for each ex post type exists in every nonempty, finite set Q , even if all agents eventually choose the same effort level.

The set of robust SPE outcomes of Γ^{FC} , which exist for both $\delta = 0$ and small $\delta > 0$, is again related to the set of solutions to an optimization problem that we call problem FC (full commitment):

$$\max_{(u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}} p_g u_{g,h} + (1 - p_g) u_{g,l} \quad (21)$$

subject to the constraints

$$p_g u_{g,h} + (1 - p_g) u_{g,l} \geq p_g u_{b,h} + (1 - p_g) u_{b,l}, \quad (22)$$

$$p_b u_{b,h} + (1 - p_b) u_{b,l} \geq p_b u_{g,h} + (1 - p_b) u_{g,l}, \quad (23)$$

$$G(\hat{d}) [p_g \Phi(u_{g,h}) + (1 - p_g) \Phi(u_{g,l})] + (1 - G(\hat{d})) [p_b \Phi(u_{b,h}) + (1 - p_b) \Phi(u_{b,l})] \leq \mathbb{E}[\tilde{y}|\hat{d}], \quad (24)$$

⁴⁶In Section 2, we have defined G on $\mathbb{R}$ such that $G(\hat{d}) = 0$ for all $\hat{d} \leq 0$.

$$\Phi(p_b u_{b,h} + (1 - p_b) u_{b,l}) \geq p_b y_h + (1 - p_b) y_l, \quad (25)$$

where $\hat{d}$ is given by

$$\hat{d} = p_g u_{g,h} + (1 - p_g) u_{g,l} - p_b u_{b,h} - (1 - p_b) u_{b,l}. \quad (26)$$

Program FC is the same as program GE($\hat{d}$) considered in section 4, with the only complication that $\hat{d}$ is now not given exogenously, but depends on contracts endogenously through (26), which captures agents' optimal effort choice in stage 2. The following lemma summarizes properties of any solution $V^{FC} = (u_{b,h}^{FC}, u_{b,l}^{FC}, u_{g,h}^{FC}, u_{g,l}^{FC})$ to problem FC, and its induced threshold for effort choice $\hat{d}^{FC} = p_g u_{g,h}^{FC} + (1 - p_g) u_{g,l}^{FC} - p_b u_{b,h}^{FC} - (1 - p_b) u_{b,l}^{FC}$.

Lemma 5. *Any solution V^{FC} to problem FC is such that $\hat{d}^{FC} \in (0, \infty)$, $u_{b,h}^{FC} = u_{b,l}^{FC} \equiv u_b^{FC}$, the constraints (23) and (24) are binding, and (22) is slack. A solution does exist.*

The lemma states that both good and bad types do exist in any solution V^{FC} , and the solutions have the same structural properties as the solutions to GE($\hat{d}$). The proof is indeed a simple extension of the proofs of Lemmas 1 and 2, but is given in Appendix B.1 for completeness. In contrast to the situation with given effort, however, the solution to FC may not be unique. The same optimal expected utility of good types may be attained by different contracts for the following reason. One solution may be such that the good types pay a high cross-subsidy to bad types, but obtain relatively much insurance, i.e. a low value of $u_{g,h} - u_{g,l}$. This leads to a low amount of available aggregate resources $\mathbb{E}[\tilde{y}|\hat{d}]$ since, at the optimum, $\hat{d}$ is given by $\hat{d} = (p_g - p_b)(u_{g,h} - u_{g,l})$ by the binding constraint (23) and by (26). Now assume the good types' incentives to invest effort are increased along their indifference curve, i.e. $u_{g,h} - u_{g,l}$ is increased without changing the good types' expected utility. Risk aversion (convexity of Φ) implies that this move requires additional resources. Resources become available, however, for two reasons. First, the reduction of good types' insurance slackens the incentive constraint (22) and makes it possible to reduce the utility given to bad types through the cross-subsidy. Second, the critical value $\hat{d}$ and thus the amount of aggregate resources is increased.⁴⁷ If these two effects together exactly compensate the increased resource requirement, an additional solution, or even a continuum of solutions, can be constructed.

Whenever there are multiple solutions to FC, they are Pareto ranked. This is because good types are equally well off in all solutions, whereas bad types are better off the higher the good types' insurance, as argued above. Moreover, since effort is chosen optimally as captured by (26), the Pareto comparison extends to ex ante utilities including effort cost.

⁴⁷While the first effect exists even if types are exogenous, the second relies on the reaction of $\hat{d}$ to the contracts. The first effect alone is thus not sufficient to generate multiple solutions.

Denote the nonempty set of solutions to FC by Ω^{FC} . Also, for any $\delta \geq 0$, let $\Omega^*(\delta)$ be the set of SPE outcomes of Γ^{FC} when the withdrawal cost parameter is δ . We then have the following result on robust equilibrium outcomes:

Proposition 5. (i) For all $\delta > 0$, $\Omega^*(\delta) \subseteq \Omega^{FC} \subseteq \Omega^*(0)$.
(ii) For each $V^{FC} \in \Omega^{FC}$ there exists a $\bar{\delta} > 0$ such that $V^{FC} \in \Omega^*(\delta)$ for all $\delta < \bar{\delta}$.

The set of equilibrium outcomes $\Omega^*(\delta)$ expands towards Ω^{FC} as δ converges to zero. Thus, in the framework with full commitment, our market model again produces Miyazaki-Wilson type contracts as the robust equilibrium outcomes, with the additional twist that effort types are endogenous. The proof is similar to the proof of Proposition 2, and can be found in Appendix B.2.

We can now use our results in order to compare robust equilibrium outcomes with and without commitment in our economy. Do equilibria without commitment lead to lower-powered incentive contracts and induce less effort, as one may expect intuitively? The following theorem provides conditions for this to be the case.

Theorem 4. Assume that condition (14) is satisfied, so that $\hat{d}^{NC}$ is unique.

(i) Suppose there exists a robust full commitment outcome $V^* \in \Omega^{FC}$ in which (25) binds. Then $V^* = V^{GE}(\hat{d}^{NC})$, i.e. it is also the robust ENC outcome in which good types exist.
(ii) For any $V^* \in \Omega^{FC}$ in which (25) does not bind, it holds that

$$\hat{d}^* > \hat{d}^{NC} \quad \text{and} \quad u_{g,h}^* - u_{g,l}^* > u_{g,h}^{GE}(\hat{d}^{NC}) - u_{g,l}^{GE}(\hat{d}^{NC}).$$

Proof. See Appendix A.5. □

The first result in Theorem 4 may be particularly surprising: If the equilibrium contracts under full commitment break even individually, they are also the unique outcome without commitment in which both types exist.⁴⁸ In that sense, for an entire class of economies, there is in fact no commitment problem at all when markets are competitive: Equilibrium contracts and incentives for effort with and without commitment fall together. The intuition is that, if the full commitment equilibrium leads to zero cross-subsidization (and hence the Rothschild-Stiglitz contracts), the composition of the population is irrelevant for firms to assess the profitability of their contract offers. In this sense, they already act as if taking the share of both types as given in the full commitment equilibrium. Then, if firms were

⁴⁸It actually follows as a corollary from the theorem that, if there is a full commitment equilibrium outcome $V^* \in \Omega^{FC}$ without cross-subsidization, it must be unique. Otherwise, if there was an additional outcome with (25) slack, this would entail less effort as argued before, and thus a strictly smaller critical value for effort choice than V^* and $V^{GE}(\hat{d}^{NC})$, contradicting statement (ii) in the theorem.

suddenly allowed to modify their contracts after effort choice, when these shares are actually fixed, their original contracts remained optimal, so that the Rothschild-Stiglitz outcome also constitutes an equilibrium without commitment.

The second result in Theorem 4 shows that the initially expected comparison between full commitment and no commitment equilibria applies if the full commitment contracts involve cross-subsidization: Aggregate effort without commitment is strictly lower, because contracts for good types are strictly lower-powered than in the full commitment outcome.⁴⁹ Yet, as we know from Section 4, even in this case the commitment problem never leads to zero effort being the unique equilibrium outcome.

6 Conclusion

We have analyzed the time-inconsistency problem with incentive contracts in a model where profit maximizing principals are competing. We have first pursued an axiomatic approach, based on weak properties that ex post market outcomes should satisfy to be considered competitive. We have shown that such outcomes Pareto dominate the allocation that a utilitarian planner can achieve, whenever they are able to sustain some incentives for effort provision. We have then provided a first example for a potentially reasonable market outcome that satisfies these requirements.

In the second part of the paper, we have provided a game-theoretic justification for the preceding results. Our game structure reflects the assumption of two-sided lack of commitment and involves a withdrawal phase to ensure equilibrium existence. Performing a robustness test based on withdrawal costs, we were able to identify robust sequential equilibria and show that they induce Miyazaki-Wilson contracts. Incentives for effort provision are preserved in these equilibria, and our general Pareto result applies. Intuitively, the market replicates a social planner who cares only about high effort agents and thus sustains maximal incentives for effort provision.

The Pareto dominance result is robust to the possibility of ex post market shutdown by the government, and also generalizes to a large class of more universal social planners. Our further results identify the level of cross-subsidization between ex post types in the market as an important property to assess outcomes. First, the comparison between market and planner becomes a strong Pareto dominance result whenever the market equilibrium entails cross-subsidization (Theorem 2). Second, even if the Pareto comparison is not applicable, because the planner does not belong to the above-mentioned class, we can still compare

⁴⁹If there are multiple equilibrium outcomes under full commitment, this holds for all of them, because all of them must involve cross-subsidization from the argument in footnote 48.

the aggregate equilibrium effort between market and planner if the market cross-subsidizes. In that case, the market performs better in terms of incentives for effort (Theorem 3). Finally, in the comparison between markets with and without commitment, a particularly strong result can be proven for the case in which the full commitment outcome involves no cross-subsidization. The market then solves the commitment problem entirely, in the sense that equilibria with and without commitment coincide (Theorem 4). Altogether, our results suggest that competitive markets are an institution that is able to deal with the commitment problem very effectively, even in a model which excludes any reputational mechanisms.

Our model provides a transparent framework in which a benevolent planner cannot replicate the outcome achieved by a market. This result is not due to exogenously assumed differences in technologies, commitment constraints, policy instruments or information, but is solely based on the different objectives that the two institutions pursue (implicitly, in case of the market). The application of this insight is not restricted to the examples that we have discussed throughout the paper (education and labor markets, insurance markets), but may be applied to banking and credit markets or even competition among rating agencies, for example. We predict that decentralized economies achieve allocations that can be Pareto superior to the outcomes under centralization, and, under even broader circumstances, will exhibit more incentive pay jobs, fewer credit defaults, and a larger per-capita social product in general.

To transparently expose the effect of competition on the commitment problem, we have ruled out reputational effects in our analysis. Future research on how competition and reputation interact may produce further interesting insights.

References

- ACEMOGLU, D., M. GOLOSOV, AND A. TSYVINSKI (2008a): “Markets versus Governments,” *Journal of Monetary Economics*, 55, 159–189.
- (2008b): “Political Economy of Mechanisms,” *forthcoming, Econometrica*.
- ALÓS-FERRER, C. (2002): “Individual Randomness in Economic Models with a Continuum of Agents,” Working paper, University of Konstanz.
- ASHEIM, G., AND T. NILSEN (1996): “Non-Discriminating Renegotiation in a Competitive Insurance Market,” *European Economic Review*, 40, 1717–1736.
- BARRO, R., AND D. GORDON (1983): “Rules, Discretion and Reputation in a Model of Monetary Policy,” *Journal of Monetary Economics*, 12, 101–121.

- BESTER, H., AND R. STRAUZ (2001): “Contracting with Imperfect Commitment and the Revelation Principle: The Single Agent Case,” *Econometrica*, 69, 1077–1098.
- BISIN, A., AND P. GOTTARDI (2006): “Efficient Competitive Equilibria with Adverse Selection,” *Journal of Political Economy*, 114, 485–516.
- BISIN, A., AND A. RAMPINI (2006): “Markets as beneficial constraints on the government,” *Journal of Public Economics*, 90, 601–629.
- BOADWAY, R., N. MARCEAU, AND M. MARCHAND (1996): “Investment in Education and the Time Inconsistency of Redistributive Tax Policy,” *Economica*, 63, 171–189.
- BUCHANAN, J. (1975): “The Samaritan’s Dilemma,” in *Altruism, Morality and Economic Theory*, ed. by E. Phelps, pp. 71–85. Russell Sage Foundation, New York.
- CHARI, V., AND P. KEHOE (1990): “Sustainable Plans,” *Journal of Political Economy*, 98, 783–802.
- CHASSAGNON, A., AND P. CHIAPPORI (1997): “Insurance under Moral Hazard and Adverse Selection: The Case of Pure Competition,” *Discussion Paper, DELTA, Paris*.
- CONCONI, P., C. PERRONI, AND R. RIEZMAN (2008): “Is partial tax harmonization desirable?,” *Journal of Public Economics*, 92, 254–267.
- DEWATRIPONT, M. (1989): “Renegotiation and Information Revelation over Time: The Case of Optimal Labor Contracts,” *Quarterly Journal of Economics*, 104, 589–620.
- DIONNE, G., AND N. DOHERTY (1994): “Adverse Selection, Commitment, and Renegotiation: Extension to and Evidence from Insurance Markets,” *Journal of Political Economy*, 102, 209–235.
- DUBEY, P., AND J. GEANAKOPOLOS (2002): “Competitive Pooling: Rothschild-Stiglitz Reconsidered,” *Quarterly Journal of Economics*, 117, 1529–1570.
- ENGERS, M., AND L. FERNANDEZ (1987): “Market Equilibrium with Hidden Knowledge and Self-Selection,” *Econometrica*, 55, 425–439.
- FARHI, E., AND I. WERNING (2008): “The Political Economy of Nonlinear Capital Taxation,” *Mimeo, Harvard University and MIT*.
- FERNANDEZ, L., AND E. RASMUSSEN (1993): “Perfectly Contestable Monopoly and Adverse Selection,” *Mimeo, Indiana University Bloomington*.

- FREIXAS, X., R. GUESNERIE, AND J. TIROLE (1985): “Planning under Incomplete Information and the Ratchet Effect,” *Review of Economic Studies*, 52, 173–191.
- FUDENBERG, D., AND J. TIROLE (1990): “Moral Hazard and Renegotiation in Agency Contracts,” *Econometrica*, 58, 1279–1320.
- GREEN, E. (1994): “Individual Level Randomness in a Nonatomic Population,” *Mimeo, University of Minnesota*.
- HART, O., A. SHLEIFER, AND R. W. VISHNY (1997): “The Proper Scope of Government: Theory and an Application to Prisons,” *Quarterly Journal of Economics*, 112, 1127–1161.
- HELLWIG, M. (1987): “Some Recent Developments in the Theory of Competition in Markets with Adverse Selection,” *European Economic Review*, 31, 319–325.
- JUDD, K. (1985): “The Law of Large Numbers with a Continuum of iid Random Variables,” *Journal of Economic Theory*, 35, 19–25.
- KEHOE, P. (1989): “Policy cooperation among benevolent governments may be undesirable,” *Review of Economic Studies*, 56, 289–296.
- KONRAD, K. (2001): “Privacy and time-consistent optimal labor income taxation,” *Journal of Public Economics*, 79, 503–519.
- KORNAI, J., E. MASKIN, AND G. ROLAND (2003): “Understanding the Soft Budget Constraint,” *Journal of Economic Literature*, 41, 1095–1136.
- KRUEGER, D., AND H. UHLIG (2006): “Competitive risk sharing contracts with one-sided commitment,” *Journal of Monetary Economics*, 53, 1661–1691.
- KUNREUTHER, H., AND M. PAULY (1985): “Market Equilibrium With Private Knowledge,” *Journal of Public Economics*, 26, 269–288.
- KYDLAND, F., AND E. PRESCOTT (1977): “Rules Rather than Discretion: The Inconsistency of Optimal Plans,” *Journal of Political Economy*, 85, 473–492.
- LA PORTA, R., F. LOPEZ-DE-SILANES, AND A. SHLEIFER (2002): “Corporate Ownership Around the World,” *Journal of Finance*, 54, 471–517.
- MIYAZAKI, H. (1977): “The Rat Race and Internal Labor Markets,” *Bell Journal of Economics*, 8, 394–418.

- NILSSEN, T. (2000): “Consumer Lock-In With Asymmetric Information,” *International Journal of Industrial Organization*, 18, 641–666.
- PHELAN, C. (1995): “Repeated Moral Hazard and One-Sided Commitment,” *Journal of Economic Theory*, 66, 488–506.
- PHELAN, C., AND E. STACCHETTI (2001): “Sequential Equilibria in a Ramsey Tax Model,” *Econometrica*, 69, 1491–1518.
- ROTHSCHILD, C. (2007): “The Efficiency of Categorical Discrimination in Insurance Markets,” *Mimeo, Middlebury College*.
- ROTHSCHILD, M., AND J. STIGLITZ (1976): “Equilibrium in Competitive Insurance Markets: An Essay in the Economics of Incomplete Information,” *Quarterly Journal of Economics*, 90, 629–649.
- SCHMIDT, K. (1996): “The Costs and Benefits of Privatization: An Incomplete Contracts Approach,” *Journal of Law, Economics and Organization*, 12, 1–24.
- (1997): “Managerial Incentives and Product Market Competition,” *Review of Economic Studies*, 64, 191–213.
- SHLEIFER, A. (1998): “State versus Private Ownership,” *Journal of Economic Perspectives*, 12, 133–150.
- STEWART, J. (1994): “The Welfare Implications of Moral Hazard and Adverse Selection in Competitive Insurance Markets,” *Economic Inquiry*, 32, 193–208.
- UHLIG, H. (1996): “A Law of Large Numbers for Large Economies,” *Economic Theory*, 8, 41–50.
- WALSH, C. E. (1995): “Optimal Contracts for Central Bankers,” *American Economic Review*, 85, 150–167.
- WILSON, C. (1977): “A Model of Insurance Markets with Incomplete Information,” *Journal of Economic Theory*, 12, 167–207.
- ZHAO, R. R. (2006): “Renegotiation-proof contract in repeated agency,” *Journal of Economic Theory*, 31, 263–281.

A Appendix

A.1 Proof of Lemma 1

(i) We first show that the statement about the constraints has to be true and that the bad types' utility will not be output-dependent in any solution $V^{SP}(\hat{d})$ to the problem, if it exists. We then prove that the problem has a unique solution. In the following, we suppress dependency on $\hat{d}$ for notational convenience.

Constraint (5). Assume that $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}$ satisfies all constraints, and (5) with slack. Consider $\tilde{V} = (u_{b,h} + \epsilon_1, u_{b,l} + \epsilon_2, u_{g,h} + \epsilon_3, u_{g,l} + \epsilon_4)$ with $\epsilon_i, i = 1, \dots, 4$, such that

$$\frac{1-p_g}{p_g}(\epsilon_2 - \epsilon_4) \leq \epsilon_3 - \epsilon_1 \leq \frac{1-p_b}{p_b}(\epsilon_2 - \epsilon_4),$$

$\epsilon_1 \geq \epsilon_2 > 0$ and $\epsilon_3 \geq \epsilon_4 > 0$. By the assumptions on $\epsilon_i, i = 1, \dots, 4$, $\tilde{V} \in \mathcal{V}$, $\tilde{V}$ satisfies (3) and (4) and $\tilde{V}$ leads to a strictly increased value of (2). To see that a set of $\epsilon_i, i = 1, \dots, 4$, with the required properties always exists, start by fixing any $\Delta_{24} \in \mathbb{R}^+$ and note that since $p_g > p_b$, there exists a $\Delta_{31} \in \mathbb{R}^+$ such that

$$\frac{1-p_g}{p_g}\Delta_{24} \leq \Delta_{31} \leq \frac{1-p_b}{p_b}\Delta_{24}.$$

Next, fix any $\epsilon_2, \epsilon_4 > 0$ such that $\epsilon_2 - \epsilon_4 = \Delta_{24}$. Clearly, it is then always possible to find $\epsilon_1, \epsilon_3 > 0$ such that $\epsilon_1 \geq \epsilon_2, \epsilon_3 \geq \epsilon_4$ and $\epsilon_3 - \epsilon_1 = \Delta_{31}$, which proves the claim. Finally, continuity of $\Phi(\cdot)$ implies that (5) is still satisfied for ϵ_i sufficiently small, so that V was not a solution to $\text{SP}(\hat{d})$.

Output-independent utilities for bad types. Assume that $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}$ with $u_{b,h} > u_{b,l}$ satisfies all constraints, and (5) with equality. Define $\tilde{u} = p_b u_{b,h} + (1-p_b)u_{b,l}$ and consider $\tilde{V} = (\tilde{u}, \tilde{u}, u_{g,h}, u_{g,l}) \in \mathcal{V}$. By construction, $\tilde{V}$ satisfies (3), and the value of (2) is the same under V and $\tilde{V}$. Since $p_g > p_b$ and $u_{b,h} > u_{b,l}$, it follows that $p_g u_{b,h} + (1-p_g)u_{b,l} > p_b u_{b,h} + (1-p_b)u_{b,l} = \tilde{u} = p_g \tilde{u} + (1-p_g)\tilde{u}$, so that $\tilde{V}$ satisfies (4) as well, given that it is satisfied by V . Strict convexity of Φ implies that $p_b \Phi(u_{b,h}) + (1-p_b)\Phi(u_{b,l}) > \Phi(p_b u_{b,h} + (1-p_b)u_{b,l}) = \Phi(\tilde{u}) = p_b \Phi(\tilde{u}) + (1-p_b)\Phi(\tilde{u})$, so that $\tilde{V}$ satisfies (5) with slack given $G(\hat{d}) \in (0, 1)$. From the previous argument, the value of the objective can then be increased above its value for $\tilde{V}$ and V , so that V was not a solution to $\text{SP}(\hat{d})$.

Constraint (4). Let $V = (u_b, u_b, u_{g,h}, u_{g,l}) \in \mathcal{V}$ satisfy all constraints, and (5) with equality. (3) and (4) together imply $u_{g,h} \geq u_b \geq u_{g,l}$. Assume (4) is slack, which implies $u_{g,h} > u_{g,l}$. Consider $\tilde{V}(\epsilon) = (u_b, u_b, u_{g,h} - \epsilon, u_{g,l} + \epsilon \frac{p_g}{1-p_g})$, $\epsilon \geq 0$, which is an element of $\mathcal{V}$ for ϵ small enough, and which satisfies $\tilde{V}(0) = V$. By construction, $\tilde{V}(\epsilon)$ satisfies (3), and the value of (2) is the same under V and $\tilde{V}(\epsilon)$, for any $\epsilon \geq 0$. (4) is also satisfied by $\tilde{V}(\epsilon)$ for ϵ small enough. Let $E_g(\epsilon) \equiv p_g \Phi(u_{g,h} - \epsilon) + (1-p_g)\Phi(u_{g,l} + \epsilon \frac{p_g}{1-p_g})$ denote the per capita expenditure for good types in $\tilde{V}(\epsilon)$. Straightforward calculations reveal that $dE_g(\epsilon)/d\epsilon < 0$ if $0 \leq \epsilon < (1-p_g)(u_{g,h} - u_{g,l})$, so that for $\epsilon > 0$ small enough, $\tilde{V}(\epsilon)$ satisfies (5) with slack given $G(\hat{d}) \in (0, 1)$. With the above argument, V cannot be a solution to $\text{SP}(\hat{d})$.

Existence and Uniqueness. The previous results show that any solution to $\text{SP}(\hat{d})$ must be of the form $V = (u_b, u_b, u_{g,h}, u_{g,l})$, and that (4) becomes $u_b = p_b u_{g,h} + (1 - p_b) u_{g,l}$, or equivalently $u_{g,l} = (u_b - p_b u_{g,h}) / (1 - p_b)$. Since $V \in \mathcal{V}$, $u_{g,h} \geq u_{g,l}$ and $p_g > p_b$ then imply that (3) is automatically satisfied. Moreover, the condition $u_{g,h} \geq u_{g,l}$ in the definition of $\mathcal{V}$ can be reformulated as $u_{g,h} \geq u_b$, or $(u_{g,h}, u_b) \in \mathcal{I}$. We can therefore state the following modified problem $\text{SP}'(\hat{d})$, which has the same solution as $\text{SP}(\hat{d})$:

$$\max_{(u_{g,h}, u_b) \in \mathcal{I}} \Psi(\hat{d}) \left[\left(\frac{1 - p_g}{1 - p_b} \right) u_b + \left(\frac{p_g - p_b}{1 - p_b} \right) u_{g,h} \right] + (1 - \Psi(\hat{d})) u_b \quad (27)$$

subject to the binding resource constraint

$$G(\hat{d}) \left[p_g \Phi(u_{g,h}) + (1 - p_g) \Phi \left(\frac{u_b - p_b u_{g,h}}{1 - p_b} \right) \right] + (1 - G(\hat{d})) \Phi(u_b) = \mathbb{E}[\tilde{y}|\hat{d}]. \quad (28)$$

Denote the LHS of (28) by $E(u_{g,h}, u_b)$. E is continuously differentiable on $\mathbb{R}^2$, and straightforward calculations reveal that it is strictly increasing in $u_{g,h}$ on $\mathcal{I}$ (including its boundary), with $\lim_{u_{g,h} \rightarrow \infty} E(u_{g,h}, u_b) = \infty$ due to convexity. E is strictly increasing in u_b globally, with $\lim_{u_b \rightarrow -\infty} E(u_{g,h}, u_b) = -\infty$.

We first claim that $u^{max} \equiv U(\mathbb{E}[\tilde{y}|\hat{d}])$ represents the largest possible choice of u_b . Consider the tuple $(u^{max}, u^{max}) \in \mathcal{I}$, which satisfies (28) by construction. Any tuple $(\tilde{u}_{g,h}, \tilde{u}_b) \in \mathcal{I}$ with $\tilde{u}_b > u^{max}$ and thus $\tilde{u}_{g,h} > u^{max}$ can be reached from (u^{max}, u^{max}) by first increasing $u_{g,h}$ from u^{max} to $\tilde{u}_{g,h}$ and then increasing u_b from u^{max} to $\tilde{u}_b$. Both moves strictly increase $E(u_{g,h}, u_b)$, so that $(\tilde{u}_{g,h}, \tilde{u}_b)$ violates (28), which proves the claim.

Now fix any $u_b \leq u^{max}$. It follows that $E(u^{max}, u_b) \leq E(u^{max}, u^{max}) = \mathbb{E}[\tilde{y}|\hat{d}]$, with strict inequality whenever $u_b < u^{max}$. Since $E(u_{g,h}, u_b)$ is strictly increasing in $u_{g,h}$ in the set $\mathcal{I}$, with $\lim_{u_{g,h} \rightarrow \infty} E(u_{g,h}, u_b) = \infty$, it follows that there exists a unique value $H(u_b)$ such that $E(H(u_b), u_b) = \mathbb{E}[\tilde{y}|\hat{d}]$, where $H(u_b) \geq u^{max} \geq u_b$. The resulting function $H : (-\infty, u^{max}] \rightarrow [u^{max}, \infty)$ is continuously differentiable and thus continuous by the implicit function theorem.

We can now reduce $\text{SP}'(\hat{d})$ to the one-dimensional problem

$$\max_{u_b \leq u^{max}} \left(1 - \Psi(\hat{d}) \frac{p_g - p_b}{1 - p_b} \right) u_b + \Psi(\hat{d}) \frac{p_g - p_b}{1 - p_b} H(u_b). \quad (29)$$

We first claim that $H(u_b)$ is strictly concave. Let $(u'_{g,h}, u'_b), (u''_{g,h}, u''_b) \in \mathcal{I}$ satisfy $E(u'_{g,h}, u'_b) = E(u''_{g,h}, u''_b) = \mathbb{E}[\tilde{y}|\hat{d}]$ and $(u'_{g,h}, u'_b) \neq (u''_{g,h}, u''_b)$. Define $u'''_{g,h} = \lambda u'_{g,h} + (1 - \lambda) u''_{g,h}$ and $u'''_b = \lambda u'_b + (1 - \lambda) u''_b$ for $\lambda \in (0, 1)$. Strict convexity of Φ then implies that $E(u'''_{g,h}, u'''_b) < \mathbb{E}[\tilde{y}|\hat{d}]$, which in turn implies that $H(u'''_b) = H(\lambda u'_b + (1 - \lambda) u''_b) > u'''_{g,h} = \lambda u'_{g,h} + (1 - \lambda) u''_{g,h} = \lambda H(u'_b) + (1 - \lambda) H(u''_b)$, which proves the claim. Second, implicit differentiation of (28) reveals that H is strictly decreasing with slope

$$H'(u_b) = \frac{G(\hat{d})(1-p_g)\Phi'(u_{g,l}) + (1-G(\hat{d}))(1-p_b)\Phi'(u_b)}{G(\hat{d})(1-p_g)p_b\Phi'(u_{g,l}) - G(\hat{d})(1-p_b)p_g\Phi'(u_{g,h})}, \quad (30)$$

where $u_{g,h} = H(u_b)$ and $u_{g,l}$ has been re-substituted for $(u_b - p_b u_{g,h})/(1-p_b)$. Observe that $\lim_{u_b \rightarrow -\infty} H'(u_b) = 0$. As u_b decreases, $u_{g,h} = H(u_b)$ increases and $u_{g,l}$ decreases. Therefore, both terms in the numerator and the first term in the denominator of (30) are decreasing as u_b is decreasing (but they remain positive). Since $\lim_{u_b \rightarrow -\infty} E(u_{g,h}, u_b) = -\infty$, it follows that $\lim_{u_b \rightarrow -\infty} H(u_b) = \infty$ and thus the second term in the denominator of (30) grows without bound as $u_b \rightarrow -\infty$, due to the Inada condition $\lim_{u \rightarrow \infty} \Phi'(u) = \infty$. Hence $\lim_{u_b \rightarrow -\infty} H'(u_b) = 0$ holds.

Strict concavity of $H(u_b)$ implies that the objective in (29) is strictly concave whenever $\Psi(\hat{d}) > 0$ and strictly increasing in u_b otherwise. Together with the fact that the objective must be strictly increasing in u_b for sufficiently small values of u_b , due to $\lim_{u_b \rightarrow -\infty} H'(u_b) = 0$ and $1 - \Psi(\hat{d})(p_g - p_b)/(1-p_b) > 0$, this implies existence and uniqueness of a solution.

(ii) We prove the claim by showing that the slope of the objective (29) evaluated at $u_b = u^{max}$ is weakly positive if $0 < \Psi(\hat{d}) \leq G(\hat{d})$. For that case, the result then follows from strict concavity of $H(u_b)$ and $H(u^{max}) = u^{max}$. The condition is

$$1 - \Psi(\hat{d}) \frac{p_g - p_b}{1 - p_b} + \Psi(\hat{d}) \frac{p_g - p_b}{1 - p_b} H'(u^{max}) \geq 0, \quad (31)$$

and after using $H(u^{max}) = u^{max}$ in (30) and some rearrangements it follows that

$$H'(u^{max}) = \frac{G(\hat{d})(p_g - p_b) - (1 - p_b)}{G(\hat{d})(p_g - p_b)}.$$

After substituting this in (31), cancelling terms and using $p_g > p_b$, we obtain that (31) is equivalent to $\Psi(\hat{d}) \leq G(\hat{d})$, which is what we assumed to start with. If $\Psi(\hat{d}) = 0$, then the objective (29) is strictly increasing in u_b and the claim follows immediately.

A.2 Proof of Lemma 2

For simplicity, we again suppress the dependency on $\hat{d}$. The argument that constraints (10) and (11) must be binding and bad types must obtain an output-independent contract is analogous to the proof of Lemma 1 (i) and therefore not repeated here.⁵⁰

Constraint (9). Let $V = (u_b, u_b, u_{g,h}, u_{g,l}) \in \mathcal{V}$ satisfy all constraints, and (9) – (11) with equality. This implies that $u_{g,h} = u_{g,l} = u_b$, and (11) simplifies to $\Phi(u_b) = \mathbb{E}[\tilde{y}|\hat{d}]$. Consider $\tilde{V}(\epsilon) = (u_b, u_b, u_b + \epsilon, u_b - \epsilon \frac{p_b}{1-p_b}) \in \mathcal{V}$ for $\epsilon \geq 0$, which satisfies $\tilde{V}(0) = V$. By construction and the fact that $p_b < p_g$, $\tilde{V}(\epsilon)$ satisfies (9) and (10), and the value of (8) is higher under $\tilde{V}(\epsilon)$ than under V for any $\epsilon > 0$. (12) is also satisfied by $\tilde{V}(\epsilon)$. Let $E(\epsilon) = G(\hat{d})[p_g \Phi(u_b + \epsilon) + (1-p_g)\Phi(u_b - \epsilon \frac{p_b}{1-p_b})] + (1-G(\hat{d}))\Phi(u_b)$ denote the average per capita expenditure in $\tilde{V}(\epsilon)$. Straightforward calculations

⁵⁰None of the arguments in the proof of Lemma 1 is affected by the additional constraint (12).

reveal that $dE(\epsilon)/d\epsilon < 0$ at $\epsilon = 0$ for $\hat{d} \in (0, \infty)$ and thus $G(\hat{d}) > 0$, so that $\tilde{V}(\epsilon)$ satisfies (11) with slack for $\epsilon > 0$ small enough. Hence V cannot be a solution to $\text{GE}(\hat{d})$.

The proof of *existence and uniqueness* of the solution to $\text{GE}(\hat{d})$ proceeds exactly as for Lemma 1 (i) to obtain a simplified optimization problem. For the special case of $\Psi(\hat{d}) = 1$ and under the additional constraint (12), which can be formulated as $u_b \geq U(p_b y_h + (1 - p_b) y_l) \equiv u^{min}$, we obtain analogously to (29),

$$u_b^{GE}(\hat{d}) = \arg \max_{u_b \in [u^{min}, u^{max}]} (1 - p_g) u_b + (p_g - p_b) H(u_b). \quad (32)$$

Existence and uniqueness now follows as for Lemma 1, with the additional simplification of a lower bound on the choice of u_b .

To show that V^{GE} satisfies conditions (C1) to (C3), consider the Rothschild-Stiglitz contracts $V^{RS} = (u_{b,h}^{RS}, u_{b,l}^{RS}, u_{g,h}^{RS}, u_{g,l}^{RS})$, which solve

$$\max_{(u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}} p_g u_{g,h} + (1 - p_g) u_{g,l}$$

subject to the constraints

$$p_k u_{k,h} + (1 - p_k) u_{k,l} \geq p_k u_{k',h} + (1 - p_k) u_{k',l} \quad \forall k, k' \in \{g, b\},$$

$$G(\hat{d}) [p_g \Phi(u_{g,h}) + (1 - p_g) \Phi(u_{g,l})] + (1 - G(\hat{d})) [p_b \Phi(u_{b,h}) + (1 - p_b) \Phi(u_{b,l})] \leq \mathbb{E}[\tilde{y}|\hat{d}],$$

$$\Phi(p_k u_{k,h} + (1 - p_k) u_{k,l}) = p_k y_h + (1 - p_k) y_l \quad \forall k, k' \in \{g, b\}.$$

Comparing with $\text{GE}(\hat{d})$, this program involves the same objective function, but a strictly smaller constraint set, implying

$$p_g u_{g,h}^{GE} + (1 - p_g) u_{g,l}^{GE} \geq p_g u_{g,h}^{RS} + (1 - p_g) u_{g,l}^{RS}.$$

Moreover, constraint (12) in $\text{GE}(\hat{d})$ implies

$$p_b u_{b,h}^{GE} + (1 - p_b) u_{b,l}^{GE} \geq U(p_b y_h + (1 - p_b) y_l) = p_b u_{b,h}^{RS} + (1 - p_b) u_{b,l}^{RS}.$$

Therefore, V^{GE} weakly Pareto-dominates V^{RS} . The result then follows from Lemma 4 in Rothschild (2007).

A.3 Proof of Proposition 2

We prove the proposition in two steps. First, we show that if $\delta > 0$, the outcome of any SPE of $\Gamma^{\hat{d}}$ must be a solution to $\text{GE}(\hat{d})$. This establishes $\Omega^*(\delta, \hat{d}) \subseteq \{V^{GE}(\hat{d})\}$ for all $\delta > 0$, the first part of statement (i). We then show that there exists a critical value $\bar{\delta} > 0$ such that $V^{GE}(\hat{d}) \in \Omega^*(\delta, \hat{d})$

for all $\delta < \bar{\delta}$, including $\delta = 0$. This establishes statement (ii), and also $\{V^{GE}(\hat{d})\} \subseteq \Omega^*(0, \hat{d})$, the second part of statement (i).

Throughout the proof, we adopt the following notation. First, we omit the asterisk indicating equilibrium strategies, for notational simplicity. Second, although σ formally is a profile of mixed strategies, we write e.g. $\sigma_j^2(s^1) = NW$ or $\sigma_j^1 = \emptyset$ to indicate a lottery placing probability 1 on a pure action. Finally, we write $\sigma = (\sigma_j, \sigma_{-j})$, $\sigma^2(s^1) = (\sigma_j^2(s^1), \sigma_{-j}^2(s^1))$ and so on. Dependence on $\hat{d}$ is suppressed for notational convenience whenever appropriate.

Step 1. Fix a value of withdrawal cost $\delta > 0$ and consider an SPE σ with outcome V^* . Observe first that $\sigma_j^2(s^1) = NW \ \forall j \in \mathcal{J}$, where s^1 is the history induced by σ , i.e. the profile of stage 2a offers. Otherwise, $\Pi_j(\sigma) = -\delta < 0$ for some $j \in \mathcal{J}$, and deviating to $\tilde{\sigma}_j^1 = \emptyset$ would be profitable. Observe also that $\Pi_j(\sigma) = 0$ for at least one $j \in \mathcal{J} \setminus \{0\}$. Otherwise, if $\Pi_j(\sigma) > 0 \ \forall j \in \mathcal{J} \setminus \{0\}$, any one of them, say i , could deviate to offering the contracts $(u_{b,h}^* + \epsilon, u_{b,l}^* + \epsilon)$ and $(u_{g,h}^* + \epsilon, u_{g,l}^* + \epsilon)$ in stage 1, for small $\epsilon > 0$, and remain active after the deviation. Since $\sigma_j^2(s^1) = NW \ \forall j \in \mathcal{J}$, the contracts available in addition to $(u_{b,h}^* + \epsilon, u_{b,l}^* + \epsilon)$ and $(u_{g,h}^* + \epsilon, u_{g,l}^* + \epsilon)$, at the end of stage 2b after the deviation, are at most those available in the SPE,⁵¹ and all agents will choose one of the deviation contracts. Also, the deviation contracts are incentive compatible, so that, for sufficiently small ϵ , the deviator could earn profits arbitrarily close to $\sum_{j \in \mathcal{J}} \Pi_j(\sigma) > \Pi_i(\sigma)$, a contradiction.

We now show that the outcome V^* must satisfy the constraints of $GE(\hat{d})$, and that it must maximize the objective (8).

Constraints (9) and (10). Incentive-compatibility is satisfied by definition of V^* .

Constraint (11). Assume to the contrary that V^* violates (11). Then there must be at least one firm $j \in \mathcal{J} \setminus \{0\}$ with $\sigma_j^2(s^1) = NW$ and $\Pi_j(\sigma) < 0$.⁵² $\tilde{\sigma}_j^1 = \emptyset$ would be a profitable deviation, which contradicts that V^* is an SPE outcome.

Constraint (12). Assume to the contrary that V^* violates (12), i.e. $\Phi(p_b u_{b,h}^* + (1 - p_b) u_{b,l}^*) < p_b y_h + (1 - p_b) y_l$. Let $\tilde{u} = p_b u_{b,h}^* + (1 - p_b) u_{b,l}^* + \epsilon$, $\epsilon > 0$, with ϵ sufficiently small to guarantee $\Phi(\tilde{u}) < p_b y_h + (1 - p_b) y_l$. The contract $(\tilde{u}, \tilde{u}) \in \mathcal{I}$ then satisfies $\pi_k(\tilde{u}, \tilde{u}) = p_k y_h + (1 - p_k) y_l - \Phi(\tilde{u}) > 0$, i.e. it earns strictly positive profits if a positive mass of agents chooses it. Consider a firm $i \in \mathcal{J}$ for which $\Pi_i(\sigma) = 0$, which exists as shown above, and assume it deviates to $\tilde{\sigma}_i^1 = \{(\tilde{u}, \tilde{u})\}$ and remains active thereafter. Since $\sigma_j^2(s^1) = NW \ \forall j \in \mathcal{J}$, the contracts that are available in addition to $(\tilde{u}, \tilde{u})$ at the end of stage 2b after the deviation are at most those available in the SPE. Hence the bad types will choose $(\tilde{u}, \tilde{u})$ and make the deviation strictly profitable, which contradicts that V^* is an SPE outcome.

Maximization of (8). Assume that V^* satisfies all constraints of $GE(\hat{d})$, but, to the contrary, $V^* \neq V^{GE}$. Then $p_g u_{g,h}^{GE} + (1 - p_g) u_{g,l}^{GE} > p_g u_{g,h}^* + (1 - p_g) u_{g,l}^*$. For $\epsilon > 0$ small enough, the contract $(u_{g,h}^{GE} - \epsilon, u_{g,l}^{GE} - \epsilon) \in \mathcal{I}$ then still satisfies $p_g(u_{g,h}^{GE} - \epsilon) + (1 - p_g)(u_{g,l}^{GE} - \epsilon) > p_g u_{g,h}^* + (1 - p_g) u_{g,l}^*$.

⁵¹If some non-deviating firms randomize in the stage 2b subgame reached after the deviation, this statement holds true for each possible outcome of the randomization.

⁵²Clearly, $\Pi_0(\sigma) = 0$ always holds.

Suppose a firm $i \in \mathcal{J}$ for which $\Pi_i(\sigma) = 0$ deviates to $\tilde{\sigma}_i^1 = \{(u_{g,h}^{GE} - \epsilon, u_{g,l}^{GE} - \epsilon), (u_b^{GE} - \epsilon, u_b^{GE} - \epsilon)\}$, with ϵ small enough as discussed, and remains active thereafter. The contracts that are additionally available at the end of stage 2b after the deviation are at most those available in the SPE, and thus the good types will choose $(u_{g,h}^{GE} - \epsilon, u_{g,l}^{GE} - \epsilon)$, given that $(u_{g,h}^*, u_{g,l}^*)$ was optimal before. Bad types weakly prefer $(u_b^{GE} - \epsilon, u_b^{GE} - \epsilon)$ over $(u_{g,h}^{GE} - \epsilon, u_{g,l}^{GE} - \epsilon)$, since V^{GE} satisfies (10). Therefore, all bad types either choose $(u_b^{GE} - \epsilon, u_b^{GE} - \epsilon)$ or a contract offered by some other firm $j \neq i$.

We claim that the deviation is strictly profitable.⁵³ Even if all bad types choose the contract $(u_b^{GE} - \epsilon, u_b^{GE} - \epsilon)$ in this outcome, the deviating firm i earns strictly positive profits.⁵⁴ If the bad types choose some other contract, firm i obtains only the good types and earns strictly positive profits as well.

Step 2. We now construct an SPE with outcome V^{GE} , which exists for sufficiently small values of δ , including $\delta = 0$.

In addition to the contracts in V^{GE} , consider the contract (u_b, u_b) that pays the expected output of bad types irrespective of actual output, so that $u_b = U(p_b y_h + (1 - p_b) y_l)$. Clearly, this contract is identical to (u_b^{GE}, u_b^{GE}) if constraint (12) is binding in V^{GE} , but the latter is strictly preferred by bad types to (u_b, u_b) otherwise. We now construct an SPE σ of $\Gamma^{\hat{d}}$ in which $\sigma_j^1 = \{(u_b, u_b)\}$ for $j = 1, 2$, $\sigma_j^1 = \{(u_b^{GE}, u_b^{GE}), (u_{g,h}^{GE}, u_{g,l}^{GE})\}$ for $j = 3, 4$, and $\sigma_j^1 = \emptyset \forall j \geq 5$. Denote the induced history by s^1 , and set $\sigma_j^2(s^1) = NW \forall j \in \mathcal{J}$. Whenever V^{GE} satisfies (12) with slack, all agents will then spread equally among firms $j = 3, 4$, which implies $\Pi_j(\sigma^2(s^1)|s^1) = 0 \forall j \in \mathcal{J}$. If (12) is satisfied with equality, bad types spread equally among firms $j = 1, \dots, 4$, while good types spread among firms 3 and 4 only. The fact that there is no cross-subsidization in V^{GE} again implies $\Pi_j(\sigma^2(s^1)|s^1) = 0 \forall j \in \mathcal{J}$. $\sigma_j^2(s^1) = NW$ is thus actually a best response for every firm in subgame $\Gamma^{\hat{d}}(s^1)$, for any value of $\delta \geq 0$, and the outcome of the SPE candidate is V^{GE} . Any potentially profitable deviation has to take place at stage 2a.

Fix a value of $\delta \geq 0$. The companies' strategies must form Nash equilibria in all off-equilibrium path subgames $\Gamma^{\hat{d}}(\tilde{s}^1)$, $\tilde{s}^1 \in S^1$, $\tilde{s}^1 \neq s^1$. The fact that each subgame $\Gamma^{\hat{d}}(\tilde{s}^1)$ is a finite normal form game implies that a Nash equilibrium does exist in each of them, possibly in mixed strategies. For each $\tilde{s}^1 \in S^1$, $\tilde{s}^1 \neq s^1$, let $\sigma^2(\tilde{s}^1)$ be such an equilibrium.⁵⁵ Now consider those stage 2b subgames $\Gamma^{\hat{d}}(\tilde{s}^1)$ that can be reached after a profitable unilateral deviation, i.e. for which there exists a firm $i \in \mathcal{J}$ such that s^1 and $\tilde{s}^1$ differ in the i th coordinate only, and where $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$. Let $\tilde{S}^1$ be the set of all histories that correspond to such subgames (suppressing the dependency on the chosen stage 2b equilibria $\sigma^2(\tilde{s}^1)$).

⁵³ Again, if there is randomization after the deviation, the following arguments apply to each outcome that occurs with positive probability.

⁵⁴ The contract $(u_b^{GE} - \epsilon, u_b^{GE} - \epsilon)$ might have been offered by non-deviators as well, in which case not all bad types choose firm i , but the deviation is still profitable.

⁵⁵ If there is more than one equilibrium in a subgame $\Gamma^{\hat{d}}(\tilde{s}^1)$, $\sigma^2(\tilde{s}^1)$ is just an arbitrary one of them.

Lemma 6. For each $\tilde{s}^1 \in \tilde{S}^1$, there exists a pure-strategy Nash equilibrium $\tilde{\sigma}^2(\tilde{s}^1)$ in $\Gamma^d(\tilde{s}^1)$.

If $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$, i.e. the deviation is still profitable under $\tilde{\sigma}^2(\tilde{s}^1)$, then $\tilde{\sigma}^2(\tilde{s}^1)$ satisfies that

- (i) each non-deviator $j \neq i, j \in \{1, 2\}$ plays $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, and
- (ii) each non-deviator $j \neq i, j \in \{3, 4\}$ plays $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$ in case of indifference, i.e. if $\Pi_j(NW, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$.

Proof. We prove the lemma by constructing the equilibrium $\tilde{\sigma}^2(\tilde{s}^1)$ from $\sigma^2(\tilde{s}^1)$.

Consider first the case where $\delta > 0$. In the given equilibrium $\sigma^2(\tilde{s}^1)$, both the deviator i and all non-deviators $j \neq i, j \in \{1, 2\}$ remain active (with probability one). For the deviator, this is because $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$ by assumption. Given that the contract (u_b, u_b) always earns non-negative profits, for non-deviators among $j \in \{1, 2\}$ remaining active even dominates withdrawal strictly. The same holds for firms $j \neq i, j \in \{3, 4\}$ if (12) is satisfied with equality in V^{GE} , because incentive compatibility and lack of cross-subsidization in V^{GE} then always implies zero profits when remaining active. Hence in that case $\sigma^2(\tilde{s}^1)$ is already in pure strategies, satisfies property (i), and (ii) is empty, so we have $\tilde{\sigma}^2(\tilde{s}^1) = \sigma^2(\tilde{s}^1)$. If (12) is slack in V^{GE} , but $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \neq \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for each non-deviator $j \neq i, j \in \{3, 4\}$, property (ii) is also empty and $\sigma^2(\tilde{s}^1)$ is in pure strategies, such that we also have $\tilde{\sigma}^2(\tilde{s}^1) = \sigma^2(\tilde{s}^1)$.

Consider then the case that (12) is slack in V^{GE} and $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for at least one $j \neq i, j \in \{3, 4\}$. Assume first that $i \notin \{3, 4\}$ in $\tilde{s}^1$. Let β_1 be the (non-random) payoff that one of firms $j \in \{3, 4\}$ would obtain if it remained active while the other did not remain active, and all other firms' strategies were as in $\sigma^2(\tilde{s}^1)$, hence pure. Let β_2 be the analogous payoff if both $j \in \{3, 4\}$ remained active, again keeping all other strategies from $\sigma^2(\tilde{s}^1)$. Indifference of (at least) one firm $j \in \{3, 4\}$ in $\sigma^2(\tilde{s}^1)$ implies that $-\delta = q\beta_1 + (1 - q)\beta_2$, where $q \in [0, 1]$ is the probability in $\sigma^2(\tilde{s}^1)$ that the other one withdraws. It must therefore be the case that either $\beta_1 < 0$ or $\beta_2 < 0$ or both. This happens if and only if the active firm(s) among 3 and 4 obtain bad types in (u_b^{GE}, u_b^{GE}) , which requires subsidization, but not enough good types in $(u_{g,h}^{GE}, u_{g,l}^{GE})$ to break even. Also, since (u_b^{GE}, u_b^{GE}) is strictly preferred to (u_b, u_b) by bad types in the present case, firms 1 and 2 do not obtain agents whenever at least one of firms 3 and 4 is active. Hence losses for active firms $j \in \{3, 4\}$ occur only if the deviator has offered a contract which is chosen by (some) good types, in the presence of $(u_{g,h}^{GE}, u_{g,l}^{GE})$, while (u_b^{GE}, u_b^{GE}) is still the best contract for bad types.

We can now distinguish two cases: first, the deviator i 's best contract for good types in $\tilde{s}^1$ could be $(u_{g,h}^{GE}, u_{g,l}^{GE})$. In this case, the deviator did not also offer (u_b^{GE}, u_b^{GE}) in $\tilde{s}^1$, because this would imply $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) = 0$ (irrespective of $\sigma_j^2(\tilde{s}^1)$, $j = 3, 4$). Hence whenever one or both firms $j \in \{3, 4\}$ are active, all bad types move only to them,⁵⁶ while all good types spread equally between them and the deviator. The number of good types that active firms $j \in \{3, 4\}$ obtain is not large enough to break even, irrespective of whether one or both of them are active, which implies $\beta_1 < 0$ and $\beta_2 < 0$. It is also straightforward to show that $\beta_1 < \beta_2$, i.e. the individual losses are smaller

⁵⁶Even if the deviator has offered an output-dependent incentive contract that leaves bad types indifferent to (u_b^{GE}, u_b^{GE}) , no bad type will choose it due to our tie-breaking assumptions.

if both $j = 3, 4$ are active and share the losses. The second possible case is that the deviator i has offered a contract in $\tilde{s}^1$ which is strictly preferred to $(u_{g,h}^{GE}, u_{g,l}^{GE})$ by good types.⁵⁷ The active firm(s) $j \in \{3, 4\}$ then obtain only the bad types and earn strictly negative profits, irrespective of whether one or both of them are active. The losses are again smaller if they are shared, also implying $\beta_1 < \beta_2 < 0$.

With these results, we can construct $\tilde{\sigma}^2(\tilde{s}^1)$ from $\sigma^2(\tilde{s}^1)$, under the assumption that $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for at least one $j \neq i, j \in \{3, 4\}$. If $i \in \{3, 4\}$, set $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, and $\tilde{\sigma}_k^2(\tilde{s}^1) = \sigma_k^2(\tilde{s}^1) \forall k \in \mathcal{J}, k \neq j$. This simply amounts to choosing an alternative best response for the indifferent player, keeping the strategies of all others. If $i \notin \{3, 4\}$, set $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$ for both $j = 3, 4$, and again $\tilde{\sigma}_k^2(\tilde{s}^1) = \sigma_k^2(\tilde{s}^1) \forall k \in \mathcal{J}, k \notin \{3, 4\}$. The fact that $\beta_1 < \beta_2 < 0$ always holds, as shown above, together with $-\delta = q\beta_1 + (1-q)\beta_2$ for a given $q \in [0, 1]$ implies $\beta_2 \geq -\delta$. The individual profits of firms $j = 3, 4$ when jointly remaining active (β_2), still given all other players' strategies from $\sigma^2(\tilde{s}^1)$, are weakly larger than $-\delta$, making it indeed a best reply to remain active. If $\tilde{\sigma}_i^2(\tilde{s}^1) = NW$ is now still a best response for the deviator, we have arrived at the desired equilibrium, because $\tilde{\sigma}^2(\tilde{s}^1)$ is a pure strategy Nash equilibrium in which all firms $j \neq i, j \in \{1, \dots, 4\}$ remain active. If i 's unique best response is now withdrawal, set $\tilde{\sigma}_i^2(\tilde{s}^1) = W$ to arrive at the final $\tilde{\sigma}^2(\tilde{s}^1)$. It is a Nash equilibrium because $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, as constructed above, is the unique best response for firms $j \neq i, j \in \{1, \dots, 4\}$ if the deviator withdraws. It is in pure strategies by construction, and properties (i) and (ii) are empty due to $\Pi_i(\tilde{\sigma}^2(\tilde{s}^1)|\tilde{s}^1) = -\delta < 0$.

Assume now that $\delta = 0$. Construct $\tilde{\sigma}^2(\tilde{s}^1)$ from $\sigma^2(\tilde{s}^1)$ by first assuming that all $j \neq i, j \in \{1, 2\}$ play $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, which is always a best response for them, and initially keep all other players' strategies as in $\sigma^2(\tilde{s}^1)$. Even if this constitutes a change of strategy from $\sigma^2(\tilde{s}^1)$, the optimal behavior of non-deviators $j \neq i, j \in \{3, 4\}$ is clearly unaffected. If $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \neq \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for all $j \neq i, j \in \{3, 4\}$, indeed keep $\tilde{\sigma}_j^2(\tilde{s}^1) = \sigma_j^2(\tilde{s}^1)$ for them. Otherwise, if $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for at least one $j \neq i, j \in \{3, 4\}$, set $\tilde{\sigma}_j^2(\tilde{s}^1) = NW \forall j \neq i, j \in \{3, 4\}$. A similar argument as for the case $\delta > 0$ implies that they then give best responses against the profile constructed so far. If $\tilde{\sigma}_i^2(\tilde{s}^1) = NW$ is still a best response of the deviator, we have arrived at the desired equilibrium. Clearly, $\tilde{\sigma}^2(\tilde{s}^1)$ is in pure strategies, it has firms $j \neq i, j \in \{1, 2\}$ remaining active, and for any firm $j \neq i, j \in \{3, 4\}$ we can have $\tilde{\sigma}^2(\tilde{s}^1) = W$ only if $\Pi_j(NW, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1) < \Pi_j(W, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$, i.e. if there is no indifference. If, on the other hand, withdrawal is now the unique best-response of the deviator, setting $\tilde{\sigma}_i^2(\tilde{s}^1) = W$ yields the desired equilibrium, because if the deviator withdraws and $\delta = 0$, all firms $j \neq i, j \in \{1, \dots, 4\}$ are indifferent between withdrawing and remaining active, making the above constructed pure strategies best responses. Furthermore, the fact that $\Pi_i(\tilde{\sigma}^2(\tilde{s}^1)|\tilde{s}^1) = -\delta = 0$ implies that (i) and (ii) are empty. $\square$

For each $\tilde{s}^1 \in \tilde{S}^1$, replace the original Nash equilibrium $\sigma^2(\tilde{s}^1)$ with the pure-strategy equilibrium

⁵⁷Any contract which leaves the good types indifferent to $(u_{g,h}^{GE}, u_{g,l}^{GE})$ but is still chosen in the presence of $(u_{g,h}^{GE}, u_{g,l}^{GE})$, must be less high-powered and would violate incentive compatibility, given that (u_b^{GE}, u_b^{GE}) is still the best contract for bad types by assumption.

$\tilde{\sigma}^2(\tilde{s}^1)$.⁵⁸ In some of the corresponding subgames, using $\tilde{\sigma}^2(\tilde{s}^1)$ might already make the deviation unprofitable, i.e. $\Pi_i(\tilde{\sigma}^2(\tilde{s}^1)|\tilde{s}^1) \leq 0$. In fact, we show in the following that this is true in all $\Gamma^{\hat{d}}(\tilde{s}^1)$, $\tilde{s}^1 \in \tilde{S}^1$, if δ is sufficiently small. To prove this claim, we assume to the contrary that there are still profitable deviations. The stage 2b equilibria reached after these deviations do then satisfy the properties (i) and (ii) of Lemma 6. To save on notation, relabel the newly constructed stage 2b equilibria back to $\sigma^2(\tilde{s}^1)$, for all $\tilde{s}^1 \in \tilde{S}^1$, and, as before, let $\tilde{S}^1$ be the set of histories that still correspond to profitable unilateral deviations from s^1 by some firm $i \in \mathcal{J}$. For each $\tilde{s}^1 \in \tilde{S}^1$, denote by $\tilde{V}(\tilde{s}^1)$ the corresponding outcome in subgame $\Gamma^{\hat{d}}(\tilde{s}^1)$, i.e. the quadruple representing the two ex post types' choices among the available contracts at the end of stage 2. $\tilde{V}(\tilde{s}^1)$ is well-defined because $\sigma^2(\tilde{s}^1)$ is in pure strategies.

Lemma 7. *There exists a value $\bar{\delta} > 0$ such that, if $0 \leq \delta < \bar{\delta}$, all outcomes $\tilde{V}(\tilde{s}^1)$, $\tilde{s}^1 \in \tilde{S}^1$, satisfy the constraints of $GE(\hat{d})$.*

Proof. Consider any $\tilde{s}^1 \in \tilde{S}^1$. By definition of $\tilde{V}(\tilde{s}^1)$ as being the outcome in $\Gamma^{\hat{d}}(\tilde{s}^1)$ under $\sigma^2(\tilde{s}^1)$, it satisfies constraints (9) and (10). (12) must also be satisfied, because the offer (u_b, u_b) remains active by construction of $\sigma^2(\tilde{s}^1)$.

Concerning (11), assume to the contrary that for some $\tilde{s}^1 \in \tilde{S}^1$, $\tilde{V}(\tilde{s}^1)$ violates (11), and let $\hat{S}^1 \subseteq \tilde{S}^1$ be the set of all such histories. As argued before, this implies losses for at least one active firm in $\Gamma^{\hat{d}}(\tilde{s}^1)$. Then, for each $\tilde{s}^1 \in \hat{S}^1$, let $\pi(\tilde{s}^1)$ be the (negative) profits of the active firm with the largest losses in $\Gamma^{\hat{d}}(\tilde{s}^1)$. We are going to show that there exists a value $\bar{\delta} > 0$ such that $\pi(\tilde{s}^1) \leq -\bar{\delta}$ for all $\tilde{s}^1 \in \hat{S}^1$, i.e. these losses are strictly bounded away from zero across all the histories $\tilde{s}^1 \in \hat{S}^1$.

Consider any $\tilde{s}^1 \in \hat{S}^1$. By assumption, $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$, and the non-deviators $j \neq i$, $j \in \{1, 2\}$ choose $\sigma_j^2(\tilde{s}^1) = NW$ and earn $\Pi_j(\sigma^2(\tilde{s}^1)|\tilde{s}^1) = 0$. Thus it must hold that V^{GE} satisfies (12) with slack and for at least one $j \neq i$, $j \in \{3, 4\}$, $\sigma_j^2(\tilde{s}^1) = NW$ and $\Pi_j(\sigma^2(\tilde{s}^1)|\tilde{s}^1) < 0$ must hold. As shown in the proof of Lemma 6, there are two cases in which this can happen. First, the deviator i 's best contract for good types in $\tilde{s}^1$ could be $(u_{g,h}^{GE}, u_{g,l}^{GE})$ and he does not offer a contract that is chosen by bad types in the presence of (u_b^{GE}, u_b^{GE}) . Denote by $\hat{S}_1^1 \subset \hat{S}^1$ the set of deviation histories with this property. Second, the deviator's best contract for good types could be strictly preferred to $(u_{g,h}^{GE}, u_{g,l}^{GE})$ by good types. Let $\hat{S}_2^1 \subset \hat{S}^1$ be the set of histories in which this is the case. Hence $\hat{S}_1^1$ and $\hat{S}_2^1$ form a partition of $\hat{S}^1$.

Consider first a history $\tilde{s}^1 \in \hat{S}_1^1$. As we have shown in the proof of Lemma 6, the profits of an active non-deviator $j \neq i$, $j \in \{3, 4\}$ are then either β_1 or β_2 , depending on whether one or both of them are active non-deviators, with $\beta_1 < \beta_2 < 0$. Hence we know that $\pi(\tilde{s}^1) \leq \max\{\beta_1, \beta_2\} = \beta_2 < 0$ for all $\tilde{s}^1 \in \hat{S}_1^1$. Consider next a history $\tilde{s}^1 \in \hat{S}_2^1$ after which active non-deviators $j \neq i$, $j \in \{3, 4\}$ obtain only bad types. They earn $\pi_h(u_b^{GE}, u_b^{GE}) < 0$ with each unit mass of bad types agents that they obtain. Given that all bad types spread equally among at most three (and thus finitely many)

⁵⁸If there are several equilibria that all satisfy the properties in Lemma 6 in a subgame $\Gamma^{\hat{d}}(\tilde{s}^1)$ for $\tilde{s}^1 \in \tilde{S}^1$, $\tilde{\sigma}^2(\tilde{s}^1)$ is just an arbitrary one of them.

firms, the losses $\pi(\tilde{s}^1)$ are strictly bounded away from zero across all $\tilde{s}^1 \in \widehat{S}_2^1$. Hence there exists a value $\beta_3 < 0$ such that $\pi(\tilde{s}^1) \leq \beta_3$ for all $\tilde{s}^1 \in \widehat{S}_2^1$.

Putting the previous results together, we obtain that $\pi(\tilde{s}^1) \leq -\bar{\delta} := \max\{\beta_2, \beta_3\} < 0$ for all $\tilde{s}^1 \in \widehat{S}^1$, i.e. whenever the outcome after a profitable deviation violates (11), a firm earns losses larger or equal to $\bar{\delta}$ in the corresponding stage 2b Nash equilibrium. But this is a contradiction if $0 \leq \delta < \bar{\delta}$, because the firm would strictly prefer to withdraw, which implies our claim. $\square$

Hence if withdrawal costs are sufficiently small, the outcome after any profitable deviation must satisfy the constraints of $\text{GE}(\hat{d})$. We next show that the outcome cannot be a solution to $\text{GE}(\hat{d})$.

Lemma 8. *If $0 \leq \delta < \bar{\delta}$, it holds that $\tilde{V}(\tilde{s}^1) \neq V^{GE}$ for all $\tilde{s}^1 \in \tilde{S}^1$.*

Proof. Assume to the contrary $\tilde{V}(\tilde{s}^1) = V^{GE}$ for some $\tilde{s}^1 \in \tilde{S}^1$. If $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$, it must be true that V^{GE} satisfies (12) with slack, the deviator i has offered $(u_{g,h}^{GE}, u_{g,l}^{GE})$ but no contract chosen by bad types in the presence of (u_b^{GE}, u_b^{GE}) , and $\sigma_j^2(\tilde{s}^1) = NW$ for at least one $j \neq i$, $j \in \{3, 4\}$. But then $\Pi_j(\sigma^2(\tilde{s}^1)|\tilde{s}^1) \leq -\bar{\delta}$, as shown in the proof of Lemma 7, which cannot occur in equilibrium if $\delta < \bar{\delta}$. $\square$

We thus know that, if $0 \leq \delta < \bar{\delta}$, after any profitable deviation history $\tilde{s}^1 \in \tilde{S}^1$ the outcome $\tilde{V}(\tilde{s}^1)$ in $\Gamma^{\hat{d}}(\tilde{s}^1)$ under $\sigma^2(\tilde{s}^1)$ must satisfy the constraints of $\text{GE}(\hat{d})$ but is not a solution to $\text{GE}(\hat{d})$. Hence good types are strictly worse off in $\tilde{V}(\tilde{s}^1)$ than in V^{GE} , which requires that $\sigma_j^2(\tilde{s}^1) = W \forall j \neq i$, $j \in \{3, 4\}$. But if some firm $j \neq i$, $j \in \{3, 4\}$ remained active instead, it would earn non-negative profits $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \geq 0$. First, it would always obtain the good types. Then, even if it obtained all bad types (in contract (u_b^{GE}, u_b^{GE})), this ensures $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \geq 0$. Hence remaining active is a best response (even unique if $\delta > 0$), contradicting that $\sigma_j^2(\tilde{s}^1) = W$, by construction of $\sigma^2(\tilde{s}^1)$. This final contradiction shows that there cannot be profitable deviations if $0 \leq \delta < \bar{\delta}$.

A.4 Proof of Lemma 3

Property (i). We will show that the solution $V^{GE}(\hat{d})$ to $\text{GE}(\hat{d})$ is continuous in $\hat{d}$ on $(0, \infty)$. It then follows that $D(\hat{d})$ is continuous as well. From the proof of Lemma 2 we know that the solution to $\text{GE}(\hat{d})$ for $\hat{d} \in (0, \infty)$ can be found by solving the simplified problem (32):

$$u_b^{GE}(\hat{d}) = \arg \max_{u_b \in [u^{min}, u^{max}(\hat{d})]} (1 - p_g)u_b + (p_g - p_b)H(u_b, \hat{d}), \quad (33)$$

where $u^{max}(\hat{d}) = U(\mathbb{E}[\tilde{y}|\hat{d}])$, and for given $\hat{d}$, the function H is continuously differentiable, strictly decreasing and strictly concave in u_b on $[u^{min}, u^{max}(\hat{d})]$.⁵⁹ Let $\mathcal{F} = (0, \infty)$, $\mathcal{U} = [u^{min}, U(p_g y_h + (1 -$

⁵⁹The dependency of $u^{max}(\hat{d})$ and $H(u_b, \hat{d})$ on $\hat{d}$ has been suppressed in earlier proofs.

$p_g)y_l]$, and $\mathcal{C}(\hat{d}) = [u^{min}, u^{max}(\hat{d})] \subset \mathcal{U}$. Clearly, the correspondence $\mathcal{C} : \mathcal{F} \rightrightarrows \mathcal{U}$ is compact-valued and continuous. Define $Z : \mathcal{U} \times \mathcal{F} \rightarrow \mathbb{R}$ by

$$Z(u_b, \hat{d}) = \begin{cases} H(u_b, \hat{d}) & \text{if } u_b \leq u^{max}(\hat{d}), \\ u^{max}(\hat{d}) & \text{if } u_b > u^{max}(\hat{d}). \end{cases} \quad (34)$$

The function Z is continuous on $\mathcal{U} \times \mathcal{F}$, because H is continuous in $\hat{d}$ and in $u_b \in [u^{min}, u^{max}(\hat{d})]$, $H(u^{max}(\hat{d}), \hat{d}) = u^{max}(\hat{d})$ holds, and $u^{max}(\hat{d})$ is continuous in $\hat{d}$. We can now rewrite the maximization problem as

$$u_b^{GE}(\hat{d}) = \arg \max_{u_b \in \mathcal{C}(\hat{d})} (1 - p_g)u_b + (p_g - p_b)Z(u_b, \hat{d}), \quad (35)$$

and Berge's maximum principle implies that $u_b^{GE}(\hat{d})$ is continuous. Then, $u_{g,h}^{GE}(\hat{d}) = Z(u_b^{GE}(\hat{d}), \hat{d})$ and $u_{g,l}^{GE}(\hat{d}) = (u_b^{GE}(\hat{d}) - p_b u_{g,h}^{GE}(\hat{d})) / (1 - p_b)$ are continuous as well.

Property (ii) Consider first the case where $\hat{d} \rightarrow 0$. We will show that, as $\hat{d} \rightarrow 0$, constraint (12) must eventually become binding in $V^{GE}(\hat{d})$, i.e. $u_b^{GE}(\hat{d}) = u^{min}$ for $\hat{d}$ small enough. Consider the slope of the objective in (33), evaluated at $u_b = u^{min}$. Using the derivative of H with respect to u_b , given in (30), the condition that the objective is weakly decreasing already in $u_b = u^{min}$ (which is then the solution to (33) due to strict concavity), can be rearranged to

$$(1 - G(\hat{d}))(p_g - p_b)\Phi'(u^{min}) \geq G(\hat{d})(1 - p_g)p_g \left[\Phi'(H(u^{min}, \hat{d})) - \Phi' \left(\frac{u^{min} - p_b H(u^{min}, \hat{d})}{1 - p_b} \right) \right]. \quad (36)$$

Fixing $u_b = u^{min}$, the budget constraint (11) can be simplified to

$$p_g \Phi(H(u^{min}, \hat{d})) + (1 - p_g) \Phi \left(\frac{u^{min} - p_b H(u^{min}, \hat{d})}{1 - p_b} \right) = p_g y_h + (1 - p_g) y_l,$$

which implies that $H(u^{min}, \hat{d})$ is independent of $\hat{d}$ and satisfies $H(u^{min}, \hat{d}) > u^{min}$. Hence the LHS of (36) converges to a strictly positive value as $\hat{d} \rightarrow 0$, while the RHS converges to zero. Hence (12) must eventually become binding, so that $\lim_{\hat{d} \rightarrow 0} u_{g,h}^{GE}(\hat{d}) = H(u^{min}, \hat{d}) > u^{min}$, and $\lim_{\hat{d} \rightarrow 0} u_{g,l}^{GE}(\hat{d}) < u^{min}$ (the latter by incentive compatibility). Hence we have $\lim_{\hat{d} \rightarrow 0} (u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d})) > 0$, which implies that $\lim_{\hat{d} \rightarrow 0} D(\hat{d}) > 0$.

Consider now the case where $\hat{d} \rightarrow \infty$. From the same arguments as above it follows that (12) must become slack for sufficiently large value of $\hat{d}$, because (36) will eventually be violated. Observe also that $u_b = u^{max}(\hat{d})$ can never be a solution to (33), for any $\hat{d} \in (0, \infty)$, as this would imply $u_{g,h} = H(u^{max}(\hat{d}), \hat{d}) = u^{max}(\hat{d}) = u_b$ and $u_{g,l} = u_b$, contradicting that $V^{GE}(\hat{d})$ satisfies (9) with slack according to Lemma 2. Hence (33) must have an interior solution for large enough $\hat{d}$. Again using the derivative of H with respect to u_b from (30), the necessary and sufficient first order

condition for (33) can then be rearranged to

$$\frac{G(\hat{d})}{1 - G(\hat{d})} = \frac{\Phi'(u_b^{GE}(\hat{d}))}{\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))} \frac{p_g - p_b}{p_g(1 - p_g)}. \quad (37)$$

Clearly, $u_b^{GE}(\hat{d})$ is bounded below by u^{min} and above by $u^{max}(\hat{d}) = U(\mathbb{E}[\tilde{y}|\hat{d}])$. Since $u^{max}(\hat{d})$ itself is bounded above by $U(p_g y_h + (1 - p_g) y_l)$, it must be that $u_b^{GE}(\hat{d}) \in [U(p_b y_h + (1 - p_b) y_l), U(p_g y_h + (1 - p_g) y_l)]$ for all $\hat{d} \in (0, \infty)$. Since $\lim_{\hat{d} \rightarrow \infty} (G(\hat{d})/(1 - G(\hat{d}))) = \infty$ while $\Phi'(u_b^{GE}(\hat{d})) \in [\Phi'(U(p_b y_h + (1 - p_b) y_l)), \Phi'(U(p_g y_h + (1 - p_g) y_l))]$ for all $\hat{d} \in (0, \infty)$, we must have $\lim_{\hat{d} \rightarrow \infty} (\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))) = 0$, since otherwise the first-order condition (37) would be violated for large $\hat{d}$. Assume $\lim_{\hat{d} \rightarrow \infty} u_{g,h}^{GE}(\hat{d}) = +\infty (-\infty)$. Then, incentive compatibility (10) requires $\lim_{\hat{d} \rightarrow \infty} u_{g,l}^{GE}(\hat{d}) = -\infty (+\infty)$, and the denominator on the RHS of (37) does not go to zero. Therefore, $\lim_{\hat{d} \rightarrow \infty} (u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d})) = 0$ has to hold (because Φ' is strictly increasing), i.e. the good types' contract converges towards full coverage. This implies $\lim_{\hat{d} \rightarrow \infty} D(\hat{d}) = 0$.

Property (iii). Assume that condition (14) ($d(\Phi''(u)/\Phi'(u))/du \geq 0$) is satisfied. Observe that this implies convexity of Φ' . We will now proceed in several steps.

First, we will show that under convexity of Φ' (and thus under (14)), both $p_g u_{g,h}^{GE}(\hat{d}) + (1 - p_g) u_{g,l}^{GE}(\hat{d})$ and $u_b^{GE}(\hat{d})$ are weakly increasing in $\hat{d}$, and strictly so if (12) is slack. As for $p_g u_{g,h}^{GE}(\hat{d}) + (1 - p_g) u_{g,l}^{GE}(\hat{d})$, this holds even without convexity of Φ' . Fix a value $\hat{d}_0 \in (0, \infty)$ and let $\hat{d} = \hat{d}_0 + \delta$ for any $\delta > 0$. In $GE(\hat{d})$, only the resource constraint (11) is affected. Straightforward calculations, using the fact that $V^{GE}(\hat{d}_0)$ satisfies (11) for $\hat{d}_0$ with equality, reveal that $V^{GE}(\hat{d}_0)$ is still feasible under $\hat{d}$ iff

$$(p_g - p_b)(y_h - y_l) - \left[p_g \Phi(u_{g,h}^{GE}(\hat{d}_0)) + (1 - p_g) \Phi(u_{g,l}^{GE}(\hat{d}_0)) - \Phi(u_b^{GE}(\hat{d}_0)) \right] \geq 0, \quad (38)$$

and satisfies the budget constraint with slack iff the inequality is strict. But the binding constraint (11) can be rearranged to

$$G(\hat{d}_0) \left[p_g \Phi(u_{g,h}^{GE}(\hat{d}_0)) + (1 - p_g) \Phi(u_{g,l}^{GE}(\hat{d}_0)) - \Phi(u_b^{GE}(\hat{d}_0)) - (p_g - p_b)(y_h - y_l) \right] + \Phi(u_b^{GE}(\hat{d}_0)) = p_b y_h + (1 - p_b) y_l,$$

which together with the fact that $\Phi(u_b^{GE}(\hat{d}_0)) \geq p_b y_h + (1 - p_b) y_l$ from (12) implies that (38) is always satisfied, and as a strict inequality whenever $\Phi(u_b^{GE}(\hat{d}_0)) > p_b y_h + (1 - p_b) y_l$. In this latter case, the optimal value of the objective under $\hat{d}$ must be strictly larger than under $\hat{d}_0$, as argued in the proof of Lemma 1. Otherwise, given that the old contracts $V^{GE}(\hat{d}_0)$ are still feasible under $\hat{d}$, the optimal value of the objective cannot decrease. Now consider the bad type's utility $u_b^{GE}(\hat{d})$. If (12) is binding, it is given by $u_b^{GE}(\hat{d}) = U(p_b y_h + (1 - p_b) y_l)$ and is independent of $\hat{d}$. Assume then that (12) is slack, such that $u_b^{GE}(\hat{d})$ satisfies the first-order condition (37). To arrive at a contradiction, suppose we increase $\hat{d}$ and $u_b^{GE}(\hat{d})$ decreases weakly. The binding self-selection constraint (10) can

be rearranged to

$$p_g u_{g,h}^{GE}(\hat{d}) + (1 - p_g) u_{g,l}^{GE}(\hat{d}) - u_b^{GE}(\hat{d}) = (p_g - p_b)(u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d})).$$

Given that $p_g u_{g,h}^{GE}(\hat{d}) + (1 - p_g) u_{g,l}^{GE}(\hat{d})$ strictly increases in $\hat{d}$, as shown above, the term $u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d})$ must also be strictly increasing. If Φ' is convex, this implies that

$$\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))$$

is increasing in $\hat{d}$, given that $u_{g,h}^{GE}(\hat{d})$ and $u_{g,l}^{GE}(\hat{d})$ cannot both decrease. Collecting results, we have that, by assumption, $u_b^{GE}(\hat{d})$ and hence $\Phi'(u_b^{GE}(\hat{d}))$ weakly decreases, while $\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))$ strictly increases. But this is a contradiction to (37), as it implies that the LHS of (37) strictly increases but the RHS strictly decreases. Hence $u_b^{GE}(\hat{d})$ is strictly increasing in $\hat{d}$ whenever (12) is slack.

Second, if (12) is slack and $u_b^{GE}(\hat{d})$ is strictly increasing at some level of $\hat{d}$, the same clearly holds for all $\hat{d}' > \hat{d}$. Together with the previous result that (12) must be binding in $V^{GE}(\hat{d})$ for sufficiently small and slack for sufficiently large values of $\hat{d}$, it follows that there exists a value $\tilde{d} \in (0, \infty)$ such that for all $\hat{d} \leq \tilde{d}$, constraint (12) will be binding in $V^{GE}(\hat{d})$ and neither $V^{GE}(\hat{d})$ nor $D(\hat{d})$ change in $\hat{d}$, while for all $\hat{d} > \tilde{d}$, (12) is slack and $u_b^{GE}(\hat{d})$ is strictly increasing in $\hat{d}$.

Third, and finally, we are going to show that, for $\hat{d} > \tilde{d}$, $u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d})$ and thus $D(\hat{d})$ are strictly decreasing in $\hat{d}$. As $\hat{d} > \tilde{d}$ grows, the LHS of the first order condition (37) grows, and so must the RHS. The condition that the derivative of the RHS of (37) with respect to $\hat{d}$ is strictly positive can be rearranged to

$$\frac{\Phi''(u_b)}{\Phi'(u_b)} > \frac{\Phi''(u_{g,h}) \frac{u'_{g,h}}{u'_b} - \Phi''(u_{g,l}) \frac{u'_{g,l}}{u'_b}}{\Phi'(u_{g,h}) - \Phi'(u_{g,l})},$$

where both the dependency on $\hat{d}$ and the superscript GE have been suppressed for notational convenience, and primes denote partial derivatives of utilities with respect to $\hat{d}$.⁶⁰ To obtain a contradiction, suppose that $u'_{g,h} \geq u'_{g,l}$. We want to find a condition under which the above inequality must be violated, that is, under which

$$\frac{\Phi''(u_b)}{\Phi'(u_b)} \leq \frac{\Phi''(u_{g,h}) \frac{u'_{g,h}}{u'_b} - \Phi''(u_{g,l}) \frac{u'_{g,l}}{u'_b}}{\Phi'(u_{g,h}) - \Phi'(u_{g,l})}. \quad (39)$$

Since $u'_b = p_b u'_{g,h} + (1 - p_b) u'_{g,l}$, we must have $u'_{g,h}/u'_b \geq 1$ and $u'_{g,l}/u'_b \leq 1$ given the assumption

⁶⁰It is straightforward to show that $u_b^{GE}(\hat{d})$, $u_{g,h}^{GE}(\hat{d})$ and $u_{g,l}^{GE}(\hat{d})$ are continuously differentiable if $\hat{d} > \tilde{d}$, given the properties of the function $H(u_b, \hat{d})$ used in (33), and the first order condition (37) for an interior solution. The derivation of the inequality makes use of the result that $u'_b > 0$.

$u'_{g,h} \geq u'_{g,l}$. Hence (39) is implied if

$$\frac{\Phi''(u_b)}{\Phi'(u_b)} \leq \frac{\Phi''(u_{g,h}) - \Phi''(u_{g,l})}{\Phi'(u_{g,h}) - \Phi'(u_{g,l})}. \quad (40)$$

This can be rearranged to

$$\Phi'(u_{g,h})\Phi'(u_b) \left[\frac{\Phi''(u_b)}{\Phi'(u_b)} - \frac{\Phi''(u_{g,h})}{\Phi'(u_{g,h})} \right] + \Phi'(u_{g,l})\Phi'(u_b) \left[\frac{\Phi''(u_{g,l})}{\Phi'(u_{g,l})} - \frac{\Phi''(u_b)}{\Phi'(u_b)} \right] \leq 0.$$

Since $u_{g,l} < u_b < u_{g,h}$, this is always satisfied under condition (14), which yields the desired contradiction. This shows that, if (14) is satisfied, $D(\hat{d})$ must be strictly decreasing in $\hat{d}$ for $\hat{d} > \tilde{d}$.

A.5 Proof of Theorem 4

Since any full commitment outcome $V^* \in \Omega^{FC}$ is a solution to FC, it must solve the corresponding necessary first-order condition. Using Lemma 5 and omitting lengthy but straightforward algebraic manipulations, we obtain

$$\begin{aligned} \frac{G(\hat{d}^*)}{1 - G(\hat{d}^*)} &= \frac{p_g - p_b}{p_g(1 - p_g)} \left[\frac{\Phi'(u_b^*)}{\Phi'(u_{g,h}^*) - \Phi'(u_{g,l}^*)} \right. \\ &\quad + \frac{g(\hat{d}^*)}{1 - G(\hat{d}^*)} \frac{\Psi(u_b^*, u_{g,h}^*, u_{g,l}^*)(1 + \mu^* \Phi'(u_b^*))}{\Phi'(u_{g,h}^*) - \Phi'(u_{g,l}^*)} \\ &\quad \left. - \frac{G(\hat{d}^*)}{1 - G(\hat{d}^*)} \frac{\mu^* \Phi'(u_b^*)}{p_g - p_b} \frac{p_g(1 - p_b)\Phi'(u_{g,h}^*) - (1 - p_g)p_b\Phi'(u_{g,l}^*)}{\Phi'(u_{g,h}^*) - \Phi'(u_{g,l}^*)} \right], \end{aligned} \quad (41)$$

where $\Psi(u_b^*, u_{g,h}^*, u_{g,l}^*) \equiv (p_g - p_b)(y_h - y_l) - [p_g\Phi(u_{g,h}^*) + (1 - p_g)\Phi(u_{g,l}^*) - \Phi(u_b^*)] \geq 0$ is the cross-subsidization from good to bad types, $\mu^* \geq 0$ is the Lagrange multiplier on constraint (25) and

$$\hat{d}^* = (p_g - p_b)(u_{g,h}^* - u_{g,l}^*). \quad (42)$$

Note also that $u_b^* = p_b u_{g,h}^* + (1 - p_b)u_{g,l}^*$ by constraint (23) and Lemma 5.

As shown in the proof of Lemma 3, for any given $\hat{d} \in (0, \infty)$, $V^{GE}(\hat{d})$ follows from the solution of the simple strictly convex maximization problem (32). We obtain the necessary and sufficient first-order condition

$$\begin{aligned} \frac{G(\hat{d})}{1 - G(\hat{d})} &= \frac{p_g - p_b}{p_g(1 - p_g)} \left[\frac{\Phi'(u_b^{GE}(\hat{d}))}{\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))} \right. \\ &\quad \left. - \frac{G(\hat{d})}{1 - G(\hat{d})} \frac{\mu(\hat{d})\Phi'(u_b^{GE}(\hat{d}))}{p_g - p_b} \frac{p_g(1 - p_b)\Phi'(u_{g,h}^{GE}(\hat{d})) - (1 - p_g)p_b\Phi'(u_{g,l}^{GE}(\hat{d}))}{\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))} \right] \end{aligned} \quad (43)$$

with $\mu(\hat{d}) \geq 0$ as Lagrange multiplier on constraint (12). In particular, this holds for the solution

$V^{GE}(\hat{d}^{NC})$ for the interior fixed point $\hat{d}^{NC} > 0$.

(i). If V^* is such that (25) binds, then $\Psi(u_b^*, u_{g,h}^*, u_{g,l}^*) = 0$. A comparison of (41) and (43) then reveals that $\hat{d}^*$, $(u_b^*, u_{g,h}^*, u_{g,l}^*)$ and μ^* also solve (43), so $V^* = V^{GE}(\hat{d}^*)$. Then, $\hat{d}^* > 0$ is a fixed point of D by (42), so that $\hat{d}^* = \hat{d}^{NC}$ and the equilibrium outcomes with and without commitment fall together.

(ii). If V^* is such that (25) does not bind, we have $\mu^* = 0$ and $\Psi(u_b^*, u_{g,h}^*, u_{g,l}^*) > 0$. (41) then implies that $\hat{d}^*$ and $(u_b^*, u_{g,h}^*, u_{g,l}^*)$ are such that

$$\frac{G(\hat{d}^*)}{1 - G(\hat{d}^*)} > \frac{p_g - p_b}{p_g(1 - p_g)} \frac{\Phi'(u_b^*)}{\Phi'(u_{g,h}^*) - \Phi'(u_{g,l}^*)} \quad (44)$$

and $\hat{d}^*$ is given by (42). On the other hand, (43) and $\mu(\hat{d}) \geq 0$ imply that, for any $\hat{d} \in (0, \infty)$,

$$\frac{G(\hat{d})}{1 - G(\hat{d})} \leq \frac{p_g - p_b}{p_g(1 - p_g)} \frac{\Phi'(u_b^{GE}(\hat{d}))}{\Phi'(u_{g,h}^{GE}(\hat{d})) - \Phi'(u_{g,l}^{GE}(\hat{d}))} \quad (45)$$

and, for the fixed point $\hat{d}^{NC} > 0$,

$$\hat{d}^{NC} = (p_g - p_b)(u_{g,h}^{GE}(\hat{d}^{NC}) - u_{g,l}^{GE}(\hat{d}^{NC})). \quad (46)$$

Fix $\bar{d}$ such that

$$(p_g - p_b)(u_{g,h}^{GE}(\bar{d}) - u_{g,l}^{GE}(\bar{d})) = \hat{d}^*. \quad (47)$$

We first show that such a value $\bar{d} \in (0, \infty)$ exists by demonstrating that

$$(p_g - p_b)(u_{g,h}^{GE}(\tilde{d}) - u_{g,l}^{GE}(\tilde{d})) > \hat{d}^*, \quad (48)$$

where $\tilde{d}$ is the critical value defined in Lemma 3. Note that $u_{g,h}^{GE}(\tilde{d}) > u_{g,l}^{GE}(\tilde{d})$,

$$p_g u_{g,h}^* + (1 - p_g) u_{g,l}^* \geq p_g u_{g,h}^{GE}(\tilde{d}) + (1 - p_g) u_{g,l}^{GE}(\tilde{d}) \quad (49)$$

by Lemma 5 and

$$p_g \Phi(u_{g,h}^*) + (1 - p_g) \Phi(u_{g,l}^*) < p_g \Phi(u_{g,h}^{GE}(\tilde{d})) + (1 - p_g) \Phi(u_{g,l}^{GE}(\tilde{d})), \quad (50)$$

since (25) does not bind in V^* . Note first that (49) and (50) rule out the two cases $u_{g,h}^* \geq u_{g,h}^{GE}(\tilde{d}) \wedge u_{g,l}^* \geq u_{g,l}^{GE}(\tilde{d})$ and $u_{g,h}^* \leq u_{g,h}^{GE}(\tilde{d}) \wedge u_{g,l}^* \leq u_{g,l}^{GE}(\tilde{d})$. Next, let us show that the case

$$u_{g,h}^* > u_{g,h}^{GE}(\tilde{d}) \wedge u_{g,l}^* < u_{g,l}^{GE}(\tilde{d}) \quad (51)$$

is also impossible. Consider all contracts $(u_{g,h}^*, u_{g,l}^*)$ that satisfy (51) and

$$p_g u_{g,h}^* + (1 - p_g) u_{g,l}^* = p_g u_{g,h}^{GE}(\tilde{d}) + (1 - p_g) u_{g,l}^{GE}(\tilde{d}).$$

Hence, we can write $u_{g,h}^* = u_{g,h}^{GE}(\tilde{d}) + \epsilon$ and $u_{g,l}^* = u_{g,l}^{GE}(\tilde{d}) - \epsilon p_g / (1 - p_g)$ with $\epsilon > 0$. Then the average cost of these contracts is

$$E_l(\epsilon) = p_g \Phi \left(u_{g,h}^{GE}(\tilde{d}) + \epsilon \right) + (1 - p_g) \Phi \left(u_{g,l}^{GE}(\tilde{d}) - \epsilon \frac{p_g}{1 - p_g} \right)$$

with $E_l'(\epsilon) > 0$ for all $\epsilon > 0$ since $u_{g,h}^{GE}(\tilde{d}) > u_{g,l}^{GE}(\tilde{d})$. Thus (50) must be violated for all contracts that satisfy (49) and (51). The only remaining case therefore is $u_{g,h}^* < u_{g,h}^{GE}(\tilde{d}) \wedge u_{g,l}^* > u_{g,l}^{GE}(\tilde{d})$, which implies $u_{g,h}^* - u_{g,l}^* < u_{g,h}^{GE}(\tilde{d}) - u_{g,l}^{GE}(\tilde{d})$ and thus (42) as claimed. Then, by Lemma 3, there exists a unique value of $\bar{d} > \tilde{d}$ such that (47) holds.

It is straightforward to see that (47) implies $u_b^* = u_b^{GE}(\bar{d})$, $u_{g,h}^* = u_{g,h}^{GE}(\bar{d})$ and $u_{g,l}^* = u_{g,l}^{GE}(\bar{d})$. This follows from the fact that $(u_b^*, u_{g,h}^*, u_{g,l}^*)$ and $(u_b^{GE}(\bar{d}), u_{g,h}^{GE}(\bar{d}), u_{g,l}^{GE}(\bar{d}))$ solve the identical resource constraint given $\hat{d}^*$, both satisfy (23) with equality and are such that $u_{g,h}^* - u_{g,l}^* = u_{g,h}^{GE}(\bar{d}) - u_{g,l}^{GE}(\bar{d})$ by (42) and (47). Then inequalities (44) and (45) imply

$$\frac{G(\bar{d})}{1 - G(\bar{d})} \leq \frac{p_g - p_b}{p_g(1 - p_g)} \frac{\Phi'(u_b^{GE}(\bar{d}))}{\Phi'(u_{g,h}^{GE}(\bar{d})) - \Phi'(u_{g,l}^{GE}(\bar{d}))} = \frac{p_g - p_b}{p_g(1 - p_g)} \frac{\Phi'(u_b^*)}{\Phi'(u_{g,h}^*) - \Phi'(u_{g,l}^*)} < \frac{G(\hat{d}^*)}{1 - G(\hat{d}^*)}$$

and therefore $\bar{d} < \hat{d}^*$. To obtain a contradiction, suppose $\hat{d}^* \leq \hat{d}^{NC}$. Then we have

$$\begin{aligned} \hat{d}^* &= (p_g - p_b)(u_{g,h}^{GE}(\bar{d}) - u_{g,l}^{GE}(\bar{d})) \\ &> (p_g - p_b)(u_{g,h}^{GE}(\hat{d}^*) - u_{g,l}^{GE}(\hat{d}^*)) \\ &\geq (p_g - p_b)(u_{g,h}^{GE}(\hat{d}^{NC}) - u_{g,l}^{GE}(\hat{d}^{NC})) = \hat{d}^{NC}, \end{aligned}$$

where the first equality follows from (47), the first inequality follows from $\tilde{d} < \bar{d} < \hat{d}^*$ and the fact that $(p_g - p_b)(u_{g,h}^{GE}(\hat{d}) - u_{g,l}^{GE}(\hat{d}))$ is strictly decreasing in $\hat{d} > \tilde{d}$ by Lemma 3 if $d[\Phi''(u)/\Phi'(u)]/du \geq 0$, the second inequality follows from the assumption that $\hat{d}^* \leq \hat{d}^{NC}$ and again Lemma 3, and the last equality follows from the fixed point condition (46). This establishes the desired contradiction, and hence $\hat{d}^* > \hat{d}^{NC}$. The result that $u_{g,h}^* - u_{g,l}^* > u_{g,h}^{GE}(\hat{d}^{NC}) - u_{g,l}^{GE}(\hat{d}^{NC})$ immediately follows from $\hat{d}^* > \hat{d}^{NC}$ and equations (42) and (46).

B Appendix (Not for Print)

B.1 Proof of Lemma 5

We first show that the statement about the constraints has to be true, that there must exist a strictly positive mass of both ex post types, and that the bad types' utility will not be output-dependent in any solution $V^{FC} = (u_{b,h}^{FC}, u_{b,l}^{FC}, u_{g,h}^{FC}, u_{g,l}^{FC})$ to the problem, if it exists. We then prove that the problem does have a solution.

Constraint (24). Assume that $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}$ satisfies all constraints, and (24) with slack. Consider $\tilde{V} = (u_{b,h} + \epsilon_1, u_{b,l} + \epsilon_2, u_{g,h} + \epsilon_3, u_{g,l} + \epsilon_4)$ with $\epsilon_i, i = 1, \dots, 4$, such that

$$\frac{1 - p_g}{p_g}(\epsilon_2 - \epsilon_4) \leq \epsilon_3 - \epsilon_1 \leq \frac{1 - p_b}{p_b}(\epsilon_2 - \epsilon_4),$$

$\epsilon_1 \geq \epsilon_2 > 0$ and $\epsilon_3 \geq \epsilon_4 > 0$. By the assumptions on $\epsilon_i, i = 1, \dots, 4$, $\tilde{V} \in \mathcal{V}$, $\tilde{V}$ satisfies (22), (23) and (25) and $\tilde{V}$ leads to a strictly increased value of (21). To see that a set of $\epsilon_i, i = 1, \dots, 4$, with the required properties always exists, start by fixing any $\Delta_{24} \in \mathbb{R}^+$ and note that since $p_g > p_b$, there exists a $\Delta_{31} \in \mathbb{R}^+$ such that

$$\frac{1 - p_g}{p_g} \Delta_{24} \leq \Delta_{31} \leq \frac{1 - p_b}{p_b} \Delta_{24}.$$

Next, fix any $\epsilon_2, \epsilon_4 > 0$ such that $\epsilon_2 - \epsilon_4 = \Delta_{24}$. Clearly, it is then always possible to find $\epsilon_1, \epsilon_3 > 0$ such that $\epsilon_1 \geq \epsilon_2$, $\epsilon_3 \geq \epsilon_4$ and $\epsilon_3 - \epsilon_1 = \Delta_{31}$, which proves the claim. Finally, continuity of $G(\cdot)$ and $\Phi(\cdot)$ implies that (24) is still satisfied for ϵ_i sufficiently small, so that V was not a solution to FC.

$\hat{d}^{FC} \in (0, \infty)$. The fact that $\hat{d}^{FC} < \infty$ in any solution immediately follows from constraints (24), (25) and our assumption that $g(d) > 0$ for all $d \geq 0$, which implies that there exist agents with arbitrarily high effort cost. Next, we show that $\hat{d}^{FC} > 0$ in any solution. Assume the contrary, i.e. suppose we have a solution $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}$ such that $\hat{d} \leq 0$. By the definition of $\hat{d}$, this implies

$$p_g u_{g,h} + (1 - p_g) u_{g,l} \leq p_b u_{b,h} + (1 - p_b) u_{b,l}. \quad (52)$$

Combining this with constraint (22) yields $p_g u_{b,h} + (1 - p_g) u_{b,l} \leq p_b u_{b,h} + (1 - p_b) u_{b,l}$ and, after cancelling and rearranging terms, $p_b(u_{b,h} - u_{b,l}) \geq p_g(u_{b,h} - u_{b,l})$. Since $p_b < p_g$ and $u_{b,h} \geq u_{b,l}$, this can only be satisfied if $u_{b,h} = u_{b,l} = u_b$ for some u_b . Hence, whenever we have a solution with $\hat{d} \leq 0$, the bad types must obtain full insurance. Substituting this in (52) and combining with constraint (22) yields $p_g u_{g,h} + (1 - p_g) u_{g,l} = u_b$, i.e. constraint (22) must be binding. Since (24) is also binding by the above result, any solution V with $\hat{d} \leq 0$ must therefore be such that $\Phi(u_b) = p_b y_h + (1 - p_b) y_l$. If $V = (u_b, u_b, u_{g,h}, u_{g,l})$ is such that $u_{g,h} > u_{g,l}$, consider first $\tilde{V} = (u_b, u_b, u_b, u_b) \in \mathcal{V}$ (otherwise, we must have $u_{g,h} = u_{g,l}$ and thus $V = \tilde{V}$ due to the binding constraint (22)). $\tilde{V}$ leaves the objective (21) unchanged compared to V since (22) is binding in V , it trivially satisfies the incentive

constraints, it satisfies constraint (25) given that V satisfies it, and it leaves (24) unaffected compared to V since $\tilde{V}$ still implies $\hat{d} = 0$. Now consider $\tilde{V}(\epsilon) = (u_b, u_b, u_b + \epsilon, u_b - \epsilon \frac{p_b}{1-p_b}) \in \mathcal{V}$ for $\epsilon \geq 0$, which satisfies $\tilde{V}(0) = \tilde{V}$. By construction and the fact that $p_b < p_g$, $\tilde{V}(\epsilon)$ satisfies (22) and (23), and the value of (21) is increased under $\tilde{V}(\epsilon)$ compared to $\tilde{V}$ (and hence V) for any $\epsilon > 0$. (25) is also satisfied by $\tilde{V}(\epsilon)$. Note that, under $\tilde{V}(\epsilon)$, $\hat{d}$ is given by

$$\hat{d}(\epsilon) = p_g \epsilon - \frac{1-p_g}{1-p_b} p_b \epsilon.$$

Observe that $\mathbb{E}[\tilde{y}|\hat{d}] = y_h - (1-p_b)(y_h - y_l) + G(\hat{d})(p_g - p_b)(y_h - y_l)$ and let

$$\begin{aligned} E(\epsilon) \equiv & G(\hat{d}(\epsilon)) \left[p_g \Phi(u_b + \epsilon) + (1-p_g) \Phi\left(u_b - \epsilon \frac{p_b}{1-p_b}\right) \right] \\ & + (1 - G(\hat{d}(\epsilon))) \Phi(u_b) - \mathbb{E}[\tilde{y}|\hat{d}] \end{aligned}$$

denote the average per capita expenditure minus resources in $\tilde{V}(\epsilon)$. Straightforward calculations reveal that

$$\left. \frac{dE(\epsilon)}{d\epsilon} \right|_{\epsilon=0} = -g(0)(p_g - p_b) \left(p_g - \frac{1-p_g}{1-p_b} p_b \right) (y_h - y_l) < 0$$

by our assumption $g(0) > 0$ and $p_b < p_g$. Hence $\tilde{V}(\epsilon)$ satisfies (24) with slack for $\epsilon > 0$ small enough and V cannot be a solution to FC.

Output-independent utilities for bad types. Assume that $V = (u_{b,h}, u_{b,l}, u_{g,h}, u_{g,l}) \in \mathcal{V}$ with $u_{b,h} > u_{b,l}$ satisfies all constraints, and (24) with equality. Define $\tilde{u} = p_b u_{b,h} + (1-p_b) u_{b,l}$ and consider $\tilde{V} = (\tilde{u}, \tilde{u}, u_{g,h}, u_{g,l}) \in \mathcal{V}$. By construction, $\tilde{V}$ satisfies (23) and (25), and the value of (21) is the same under V and $\tilde{V}$. Since $p_g > p_b$ and $u_{b,h} > u_{b,l}$, it follows that $p_g u_{b,h} + (1-p_g) u_{b,l} > p_b u_{b,h} + (1-p_b) u_{b,l} = \tilde{u} = p_g \tilde{u} + (1-p_g) \tilde{u}$, so that $\tilde{V}$ satisfies (22) as well, given that it is satisfied by V . Moreover, by (26), $\tilde{V}$ implies the same value for $\hat{d}$ as V . Strict convexity of Φ implies that $p_b \Phi(u_{b,h}) + (1-p_b) \Phi(u_{b,l}) > \Phi(p_b u_{b,h} + (1-p_b) u_{b,l}) = \Phi(\tilde{u}) = p_b \Phi(\tilde{u}) + (1-p_b) \Phi(\tilde{u})$, so that $\tilde{V}$ satisfies (24) with slack given $\hat{d} \in (0, \infty)$. From the previous argument, the value of the objective can then be increased above its value for $\tilde{V}$ and V , so that V was not a solution to FC.

Constraint (23). Let $V = (u_b, u_b, u_{g,h}, u_{g,l}) \in \mathcal{V}$ satisfy all constraints, and (24) with equality. (22) and (23) together imply $u_{g,h} \geq u_b \geq u_{g,l}$. Assume (23) is slack, which implies $u_{g,h} > u_{g,l}$. Consider $\tilde{V}(\epsilon) = (u_b, u_b, u_{g,h} - \epsilon, u_{g,l} + \epsilon \frac{p_g}{1-p_g})$, $\epsilon \geq 0$, which is an element of $\mathcal{V}$ for ϵ small enough, and which satisfies $\tilde{V}(0) = V$. By construction, $\tilde{V}(\epsilon)$ satisfies (22) and (25), and the value of (21) is the same under V and $\tilde{V}(\epsilon)$, for any $\epsilon \geq 0$. (23) is also satisfied by $\tilde{V}(\epsilon)$ for ϵ small enough. Let $E_g(\epsilon) \equiv p_g \Phi(u_{g,h} - \epsilon) + (1-p_g) \Phi(u_{g,l} + \epsilon \frac{p_g}{1-p_g})$ denote the per capita expenditure for good types in $\tilde{V}(\epsilon)$. Straightforward calculations reveal that $dE_g(\epsilon)/d\epsilon < 0$ if $0 \leq \epsilon < (1-p_g)(u_{g,h} - u_{g,l})$. Note that, for any $\epsilon \geq 0$, $\tilde{V}(\epsilon)$ implies the same $\hat{d}$ as V by (26), so that for $\epsilon > 0$ small enough, $\tilde{V}(\epsilon)$ satisfies (24) with slack given $\hat{d} \in (0, \infty)$. With the above argument, V cannot be a solution to FC.

Constraint (22). Let $V = (u_b, u_b, u_{g,h}, u_{g,l}) \in \mathcal{V}$ satisfy all constraints, and (22) – (24) with

equality. This implies that $u_{g,h} = u_{g,l} = u_b$, and (26) implies $\hat{d} = 0$. The fact that this cannot be a solution to FC then follows from the above proof that $\hat{d} > 0$ in any solution.

Existence. The previous results show that any solution to FC must be of the form $V = (u_b, u_b, u_{g,h}, u_{g,l})$, and that (22) and (23) can be replaced by $u_b = p_b u_{g,h} + (1 - p_b) u_{g,l}$, or equivalently $u_{g,l} = (u_b - p_b u_{g,h}) / (1 - p_b)$. Using this equation, the condition $u_{g,h} \geq u_{g,l}$ in the definition of $\mathcal{V}$ can be reformulated as $u_{g,h} \geq u_b$, or $(u_{g,h}, u_b) \in \mathcal{I}$. We can therefore state the following modified problem (FC'), which has the same solution as FC:

$$\max_{(u_{g,h}, u_b) \in \mathcal{I}} \left(\frac{1 - p_g}{1 - p_b} \right) u_b + \left(\frac{p_g - p_b}{1 - p_b} \right) u_{g,h} \quad (53)$$

subject to the constraints

$$\begin{aligned} G(\hat{d}) \left[p_g \Phi(u_{g,h}) + (1 - p_g) \Phi \left(\frac{u_b - p_b u_{g,h}}{1 - p_b} \right) \right] + (1 - G(\hat{d})) \Phi(u_b) \\ - G(\hat{d})(p_g - p_b)(y_h - y_l) = y_h - (1 - p_b)(y_h - y_l), \end{aligned} \quad (54)$$

$$u_b \geq U(p_b y_h + (1 - p_b) y_l), \quad (55)$$

where (54) is the binding resource constraint, (55) the cross-subsidization constraint, and $\hat{d}$ is given by

$$\hat{d} = \frac{p_g - p_b}{1 - p_b} (u_{g,h} - u_b). \quad (56)$$

We prove existence by verifying the conditions of the Weierstrass theorem. Clearly, the objective function in (53) is continuous. Denote the constraint set by $\mathcal{S}$, i.e. let $\mathcal{S}$ be the set of pairs $(u_{g,h}, u_b)$ such that $(u_{g,h}, u_b) \in \mathcal{I}$ and the constraints (54) to (56) are satisfied. $\mathcal{S}$ is closed by continuity of Φ and G and the fact that the constraints in the definition of $\mathcal{I}$ and (55) are weak inequalities.

We next prove boundedness of $\mathcal{S}$. $\mathcal{S}$ is obviously bounded below in terms of both $u_{g,h}$ and u_b by (55) and the constraint $u_{g,h} \geq u_b$ from $(u_{g,h}, u_b) \in \mathcal{I}$. Hence, all that remains to be proven is the fact that $\mathcal{S}$ is bounded above with respect to both $u_{g,h}$ and u_b . Let us start with fixing an arbitrary $u_{g,h} \geq U(p_b y_h + (1 - p_b) y_l)$. Then the set of values of u_b such that $(u_{g,h}, u_b) \in \mathcal{S}$ is bounded above since $u_b \leq u_{g,h}$ from $(u_{g,h}, u_b) \in \mathcal{I}$. Next, fix some $u_b \geq U(p_b y_h + (1 - p_b) y_l)$. The set of values of $u_{g,h}$ such that $(u_{g,h}, u_b) \in \mathcal{S}$ is bounded above since (54) must be violated for sufficiently high $u_{g,h}$. To see this, denote the square-bracketed term on the LHS of (54) by $\Lambda(u_{g,h}, u_b)$. Using the fact that $p_b < p_g$, it is straightforward to show that $\Lambda(u_{g,h}, u_b)$ is strictly increasing and strictly convex in $u_{g,h}$ for all $u_{g,h} \geq u_b$. Since $\hat{d}$ is non-decreasing in $u_{g,h}$ by (56), it must be that the LHS of (54) grows without bound as $u_{g,h} \rightarrow \infty$ for any given u_b . Since the RHS of (54) is fixed and finite, (54) must be violated for sufficiently high values of $u_{g,h}$. Finally, let both $u_{g,h} \rightarrow \infty$ and $u_b \rightarrow \infty$ with $u_{g,h} \geq u_b$. First, we have $\Phi(u_b) \rightarrow \infty$ for $u_b \rightarrow \infty$. Moreover, since $p_b < p_g$, it is straightforward to show that $\Lambda(u_{g,h}, u_b)$ is strictly increasing and strictly convex in both $u_{g,h}$ and u_b for all $(u_{g,h}, u_b) \in \mathcal{I}$. We therefore must have $\Lambda(u_{g,h}, u_b) \rightarrow \infty$ as $u_{g,h} \rightarrow \infty$ and $u_b \rightarrow \infty$ with $u_{g,h} \geq u_b$. Together

with $G(\hat{d}) \in [0, 1]$ for all $(u_{g,h}, u_b)$, this shows that the LHS of (54) grows without bound as we let $u_{g,h} \rightarrow \infty$ and $u_b \rightarrow \infty$ with $u_{g,h} \geq u_b$. Since the RHS of (54) is fixed and finite, (54) must be violated for sufficiently high $u_{g,h}$ and u_b such that $(u_{g,h}, u_b) \in \mathcal{I}$. This completes the proof of boundedness of $\mathcal{S}$, and hence existence of a solution follows from the Weierstrass theorem.

B.2 Proof of Proposition 5

As for Proposition 2, we proceed in two steps. First, we show that if $\delta > 0$, the outcome of any SPE of Γ^{FC} must be a solution to FC. This establishes $\Omega^*(\delta) \subseteq \Omega^{FC}$ for all $\delta > 0$, the first part of statement (i). We then show that, for any $V^{FC} \in \Omega^{FC}$, there exists a critical value $\bar{\delta} > 0$ such that $V^{FC} \in \Omega^*(\delta)$ for all $\delta < \bar{\delta}$, including $\delta = 0$. This establishes statement (ii), and also $\Omega^{FC} \subseteq \Omega^*(0)$, the second part of statement (i).

Step 1. Fix a value of withdrawal cost $\delta > 0$ and consider an SPE σ with outcome V^* . Observe first that $\sigma_j^2(s^1) = NW \ \forall j \in \mathcal{J}$, where s^1 is the history induced by σ , i.e. the profile of stage 1a offers. Otherwise, $\Pi_j(\sigma) = -\delta < 0$ for some $j \in \mathcal{J}$, and deviating to $\tilde{\sigma}_j^1 = \emptyset$ would be profitable. Observe also $\Pi_j(\sigma) = 0$ for at least one $j \in \mathcal{J} \setminus \{0\}$. Otherwise, if $\Pi_j(\sigma) > 0 \ \forall j \in \mathcal{J} \setminus \{0\}$, any one of them, say i , could deviate to offering the contracts $(u_{b,h}^* + \epsilon, u_{b,l}^* + \epsilon)$ and $(u_{g,h}^* + \epsilon, u_{g,l}^* + \epsilon)$ in stage 1a, for small $\epsilon > 0$, and remain active after the deviation. Since $\sigma_j^2(s^1) = NW \ \forall j \in \mathcal{J}$, the contracts available in addition to $(u_{b,h}^* + \epsilon, u_{b,l}^* + \epsilon)$ and $(u_{g,h}^* + \epsilon, u_{g,l}^* + \epsilon)$, at the end of stage 1b after the deviation, are at most those available in the SPE,⁶¹ and all agents will choose one of the deviation contracts. Also, the deviation contracts are incentive compatible, and induce the same critical value $\hat{d}^*$ as in V^* , so that, for sufficiently small ϵ , the deviator could earn profits arbitrarily close to $\sum_{j \in \mathcal{J}} \Pi_j(\sigma) > \Pi_i(\sigma)$, a contradiction.

We now show that the outcome V^* must satisfy the constraints of FC, and that it must maximize the objective (21).

Constraints (22) and (23). Incentive-compatibility is satisfied by definition of V^* .

Constraint (24). Assume to the contrary that V^* violates (24). The equilibrium critical value $\hat{d}^*$ is identical to the one given by (26) and used in (24), so that there must be at least one firm $j \in \mathcal{J} \setminus \{0\}$ with $\sigma_j^2(s^1) = NW$ and $\Pi_j(\sigma) < 0$.⁶² $\tilde{\sigma}_j^1 = \emptyset$ would be a profitable deviation, which contradicts that V^* is an SPE outcome.

Constraint (25). Assume to the contrary that V^* violates (25), i.e. $\Phi(p_b u_{b,h}^* + (1 - p_b) u_{b,l}^*) < p_b y_h + (1 - p_b) y_l$. Let $\tilde{u} = p_b u_{b,h}^* + (1 - p_b) u_{b,l}^* + \epsilon$, $\epsilon > 0$, with ϵ sufficiently small to guarantee $\Phi(\tilde{u}) < p_b y_h + (1 - p_b) y_l$. The contract $(\tilde{u}, \tilde{u}) \in \mathcal{I}$ then satisfies $\pi_b(\tilde{u}, \tilde{u}) = p_b y_h + (1 - p_b) y_l - \Phi(\tilde{u}) > 0$, i.e. it earns strictly positive profits if a positive mass of agents (who would become bad types) chooses it. Consider a firm $i \in \mathcal{J}$ for which $\Pi_i(\sigma) = 0$, which exists as shown above, and assume it deviates

⁶¹If some non-deviating firms randomize in the stage 1b subgame reached after the deviation, this statement holds true for each possible outcome of the randomization.

⁶²Clearly, $\Pi_0(\sigma) = 0$ always holds.

to $\tilde{\sigma}_i^1 = \{(\tilde{u}, \tilde{u})\}$ and remains active thereafter. Since $\sigma_j^2(s^1) = NW \ \forall j \in \mathcal{J}$, the contracts that are available in addition to $(\tilde{u}, \tilde{u})$ at the end of stage 1b after the deviation are at most those available in the SPE. Hence if some agents decide to become bad types after the deviation, they will choose $(\tilde{u}, \tilde{u})$ and make the deviation strictly profitable. But there are always bad types as argued in the proof of Lemma 5, which contradicts that V^* is an SPE outcome.

Maximization of (21). Assume that V^* satisfies all constraints of FC, but, to the contrary, $V^* \notin \Omega^{FC}$. Fix an arbitrary solution $V^{FC} \in \Omega^{FC}$. Then $p_g u_{g,h}^{FC} + (1 - p_g) u_{g,l}^{FC} > p_g u_{g,h}^* + (1 - p_g) u_{g,l}^*$. For $\epsilon > 0$ small enough, the contract $(u_{g,h}^{FC} - \epsilon, u_{g,l}^{FC} - \epsilon) \in \mathcal{I}$ then still satisfies $p_g(u_{g,h}^{FC} - \epsilon) + (1 - p_g)(u_{g,l}^{FC} - \epsilon) > p_g u_{g,h}^* + (1 - p_g) u_{g,l}^*$. Suppose a firm $i \in \mathcal{J}$ for which $\Pi_i(\sigma) = 0$ deviates to $\tilde{\sigma}_i^1 = \{(u_{g,h}^{FC} - \epsilon, u_{g,l}^{FC} - \epsilon), (u_{b,h}^{FC} - \epsilon, u_{b,l}^{FC} - \epsilon)\}$, with ϵ small enough as discussed, and remains active thereafter. The contracts that are additionally available at the end of stage 1b after the deviation are at most those available in the SPE, and thus if some agents decide to become good types they will choose $(u_{g,h}^{FC} - \epsilon, u_{g,l}^{FC} - \epsilon)$, given that $(u_{g,h}^*, u_{g,l}^*)$ was optimal before. Bad types weakly prefer $(u_{b,h}^{FC} - \epsilon, u_{b,l}^{FC} - \epsilon)$ over $(u_{g,h}^{FC} - \epsilon, u_{g,l}^{FC} - \epsilon)$, since V^{FC} satisfies (23). Therefore, all bad types either choose $(u_{b,h}^{FC} - \epsilon, u_{b,l}^{FC} - \epsilon)$ or a contract offered by some other firm $j \neq i$.

We claim that the deviation is strictly profitable if there is a positive mass of good types in the outcome after the deviation.⁶³ Even if all bad types choose the contract $(u_{b,h}^{FC} - \epsilon, u_{b,l}^{FC} - \epsilon)$ in this outcome, then by construction the critical value $\hat{d}$ is equal to $\hat{d}^{FC}$, the critical value in V^{FC} , which implies that the deviating firm i earns strictly positive profits.⁶⁴ If the bad types choose some other contract, firm i obtains only the good types and earns strictly positive profits as well. We finally show that there must actually be a positive mass of good types in the deviation outcome. First, if the bad types' optimal contract after the deviation is $(u_{b,h}^{FC} - \epsilon, u_{b,l}^{FC} - \epsilon)$, the critical value for effort choice is equal to $\hat{d}^{FC} > 0$ by Lemma 5. Otherwise, the bad types' optimal contract was already available in the SPE, so they cannot be better off than in V^* . On the other hand, potential good types are strictly better off after the deviation. Note that $\hat{d}^* = p_g u_{g,h}^* + (1 - p_g) u_{g,l}^* - p_b u_{b,h}^* - (1 - p_b) u_{b,l}^* \geq 0$ by incentive compatibility of V^* , $p_g > p_b$, and $u_{b,h}^* \geq u_{b,l}^*$. But if the good types' utility strictly increases while the bad types' utility does not increase, the critical value $\hat{d}$ after the deviation and hence the share of good types must be strictly positive.

Step 2. For each $V^{FC} \in \Omega^{FC}$, we construct an SPE with outcome V^{FC} , which exists for sufficiently small values of δ , including $\delta = 0$.

In addition to the contracts in V^{FC} , consider the contract (u_b, u_b) that pays the expected output of bad types irrespective of actual output, so that $u_b = U(p_b y_h + (1 - p_b) y_l)$. Clearly, this contract is identical to $(u_{b,h}^{FC}, u_{b,l}^{FC})$ if constraint (25) is binding in V^{FC} , but the latter is strictly preferred by bad types to (u_b, u_b) otherwise. We now construct an SPE σ of Γ^{FC} in which $\sigma_j^1 = \{(u_b, u_b)\}$

⁶³Again, if there is randomization after the deviation, the following arguments apply to each outcome that occurs with positive probability.

⁶⁴The contract $(u_{b,h}^{FC} - \epsilon, u_{b,l}^{FC} - \epsilon)$ might have been offered by non-deviators as well, in which case not all bad types choose firm i , but the statement about the critical value $\hat{d}$ and the positive profits is still true.

for $j = 1, 2$, $\sigma_j^1 = \{(u_{b,h}^{FC}, u_{b,l}^{FC}), (u_{g,h}^{FC}, u_{g,l}^{FC})\}$ for $j = 3, 4$, and $\sigma_j^1 = \emptyset \forall j \geq 5$. Denote the induced history by s^1 , and set $\sigma_j^2(s^1) = NW \forall j \in \mathcal{J}$. Whenever V^{FC} satisfies (25) with slack, all agents will then spread equally among firms $j = 3, 4$, which implies the critical value $\hat{d}^{FC}$ and $\Pi_j(\sigma^2(s^1)|s^1) = 0 \forall j \in \mathcal{J}$. If (25) is satisfied with equality, the critical value is still $\hat{d}^{FC}$, bad types spread equally among firms $j = 1, \dots, 4$, while good types spread among firms 3 and 4 only. The fact that there is no cross-subsidization in V^{FC} again implies $\Pi_j(\sigma^2(s^1)|s^1) = 0 \forall j \in \mathcal{J}$. $\sigma_j^2(s^1) = NW$ is thus actually a best response for every firm in subgame $\Gamma^{FC}(s^1)$, for any value of $\delta \geq 0$, and the outcome of the SPE candidate is V^{FC} . Any potentially profitable deviation has to take place at stage 1a.

Fix a value of $\delta \geq 0$. The companies' strategies must form Nash equilibria in all off-equilibrium path subgames $\Gamma^{FC}(\tilde{s}^1)$, $\tilde{s}^1 \in S^1$, $\tilde{s}^1 \neq s^1$. The fact that each subgame $\Gamma^{FC}(\tilde{s}^1)$ is a finite normal form game implies that a Nash equilibrium does exist in each of them, possibly in mixed strategies. For each $\tilde{s}^1 \in S^1$, $\tilde{s}^1 \neq s^1$, let $\sigma^2(\tilde{s}^1)$ be such an equilibrium.⁶⁵ Now consider those stage 1b subgames $\Gamma^{FC}(\tilde{s}^1)$ that can be reached after a profitable unilateral deviation, i.e. for which there exists a firm $i \in \mathcal{J}$ such that s^1 and $\tilde{s}^1$ differ in the i th coordinate only, and where $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$. Let $\tilde{S}^1$ be the set of all histories that correspond to such subgames (suppressing the dependency on the chosen stage 1b equilibria $\sigma^2(\tilde{s}^1)$).

Lemma 9. *For each $\tilde{s}^1 \in \tilde{S}^1$, there exists a pure-strategy Nash equilibrium $\tilde{\sigma}^2(\tilde{s}^1)$ in $\Gamma^{FC}(\tilde{s}^1)$. If $\Pi_i(\tilde{\sigma}^2(\tilde{s}^1)|\tilde{s}^1) > 0$, i.e. the deviation is still profitable under $\tilde{\sigma}^2(\tilde{s}^1)$, then $\tilde{\sigma}^2(\tilde{s}^1)$ satisfies that*

- (i) *each non-deviator $j \neq i$, $j \in \{1, 2\}$ plays $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, and*
- (ii) *each non-deviator $j \neq i$, $j \in \{3, 4\}$ plays $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$ in case of indifference, i.e. if $\Pi_j(NW, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$.*

Proof. We prove the lemma by constructing the equilibrium $\tilde{\sigma}^2(\tilde{s}^1)$ from $\sigma^2(\tilde{s}^1)$.

Consider first the case where $\delta > 0$. In the given equilibrium $\sigma^2(\tilde{s}^1)$, both the deviator i and all non-deviators $j \neq i$, $j \in \{1, 2\}$ remain active (with probability one). For the deviator, this is because $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$ by assumption. Given that the contract (u_b, u_b) always earns zero profits (agents who choose it become bad types), for non-deviators among $j \in \{1, 2\}$ remaining active even dominates withdrawal strictly. The same holds for firms $j \neq i$, $j \in \{3, 4\}$ if (25) is satisfied with equality in V^{FC} , because incentive compatibility and lack of cross-subsidization in V^{FC} then always implies zero profits when remaining active. Hence in that case $\sigma^2(\tilde{s}^1)$ is already in pure strategies, satisfies property (i), and (ii) is empty, so we have $\tilde{\sigma}^2(\tilde{s}^1) = \sigma^2(\tilde{s}^1)$. If (25) is slack in V^{FC} , but $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \neq \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for each non-deviator $j \neq i$, $j \in \{3, 4\}$, property (ii) is also empty and $\sigma^2(\tilde{s}^1)$ is in pure strategies, such that we also have $\tilde{\sigma}^2(\tilde{s}^1) = \sigma^2(\tilde{s}^1)$.

Consider then the case that (25) is slack in V^{FC} and $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for at least one $j \neq i$, $j \in \{3, 4\}$. Assume first that $i \notin \{3, 4\}$ in $\tilde{s}^1$. Let β_1 be the (non-random) payoff that one of firms $j \in \{3, 4\}$ would obtain if it remained active while the other did not remain active, and all other firms' strategies were as in $\sigma^2(\tilde{s}^1)$, hence pure. Let β_2 be the analogous payoff

⁶⁵If there is more than one equilibrium in a subgame $\Gamma^{FC}(\tilde{s}^1)$, $\sigma^2(\tilde{s}^1)$ is just an arbitrary one of them.

if both $j \in \{3, 4\}$ remained active, again keeping all other strategies from $\sigma^2(\tilde{s}^1)$. Indifference of (at least) one firm $j \in \{3, 4\}$ in $\sigma^2(\tilde{s}^1)$ implies that $-\delta = q\beta_1 + (1 - q)\beta_2$, where $q \in [0, 1]$ is the probability in $\sigma^2(\tilde{s}^1)$ that the other one withdraws. It must therefore be the case that either $\beta_1 < 0$ or $\beta_2 < 0$ or both. This happens if and only if the active firm(s) among 3 and 4 obtain bad types in $(u_{b,h}^{FC}, u_{b,l}^{FC})$, which requires subsidization, but not enough good types in $(u_{g,h}^{FC}, u_{g,l}^{FC})$ to break even. Observe that the induced critical value $\hat{d}$ is the same irrespective of whether one or both of firms 3 and 4 remain active. Also, since $(u_{b,h}^{FC}, u_{b,l}^{FC})$ is strictly preferred to (u_b, u_b) by bad types in the present case, firms 1 and 2 do not obtain agents whenever at least one of firms 3 and 4 is active. Hence losses for active firms $j \in \{3, 4\}$ occur only if the deviator has offered a contract which is chosen by (some) good types, in the presence of $(u_{g,h}^{FC}, u_{g,l}^{FC})$, while $(u_{b,h}^{FC}, u_{b,l}^{FC})$ is still the best contract for bad types.

We can now distinguish two cases: first, the deviator i 's best contract for good types in $\tilde{s}^1$ could be $(u_{g,h}^{FC}, u_{g,l}^{FC})$. In this case, the deviator did not also offer $(u_{b,h}^{FC}, u_{b,l}^{FC})$ in $\tilde{s}^1$, because this would imply a critical value $\hat{d}^{FC}$ and $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) = 0$ (irrespective of $\sigma_j^2(\tilde{s}^1)$, $j = 3, 4$). Hence whenever one or both firms $j \in \{3, 4\}$ are active, all bad types move only to them,⁶⁶ while all good types spread equally between them and the deviator. The induced critical value is $\hat{d}^{FC}$, and the number of good types that active firms $j \in \{3, 4\}$ obtain is not large enough to break even, irrespective of whether one or both of them are active, which implies $\beta_1 < 0$ and $\beta_2 < 0$. It is also straightforward to show that $\beta_1 < \beta_2$, i.e. the individual losses are smaller if both $j = 3, 4$ are active and share the losses. The second possible case is that the deviator i has offered a contract in $\tilde{s}^1$ which is strictly preferred to $(u_{g,h}^{FC}, u_{g,l}^{FC})$ by good types.⁶⁷ The active firm(s) $j \in \{3, 4\}$ then obtain only bad types (which always exist as argued before) and earn strictly negative profits, irrespective of whether one or both of them are active. The losses are again smaller if they are shared, also implying $\beta_1 < \beta_2 < 0$.

With these results, we can construct $\tilde{\sigma}^2(\tilde{s}^1)$ from $\sigma^2(\tilde{s}^1)$, under the assumption that $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for at least one $j \neq i$, $j \in \{3, 4\}$. If $i \in \{3, 4\}$, set $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, and $\tilde{\sigma}_k^2(\tilde{s}^1) = \sigma_k^2(\tilde{s}^1) \forall k \in \mathcal{J}, k \neq j$. This simply amounts to choosing an alternative best response for the indifferent player, keeping the strategies of all others. If $i \notin \{3, 4\}$, set $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$ for both $j = 3, 4$, and again $\tilde{\sigma}_k^2(\tilde{s}^1) = \sigma_k^2(\tilde{s}^1) \forall k \in \mathcal{J}, k \notin \{3, 4\}$. The fact that $\beta_1 < \beta_2 < 0$ always holds, as shown above, together with $-\delta = q\beta_1 + (1 - q)\beta_2$ for a given $q \in [0, 1]$ implies $\beta_2 \geq -\delta$. The individual profits of firms $j = 3, 4$ when jointly remaining active (β_2), still given all other players' strategies from $\sigma^2(\tilde{s}^1)$, are weakly larger than $-\delta$, making it indeed a best reply to remain active. If $\tilde{\sigma}_i^2(\tilde{s}^1) = NW$ is now still a best response for the deviator, we have arrived at the desired equilibrium, because $\tilde{\sigma}^2(\tilde{s}^1)$ is a pure strategy Nash equilibrium in which all firms $j \neq i, j \in \{1, \dots, 4\}$

⁶⁶Even if the deviator has offered an output-dependent incentive contract that leaves bad types indifferent to $(u_{b,h}^{FC}, u_{b,l}^{FC})$, no bad type will choose it due to our tie-breaking assumptions.

⁶⁷Any contract which leaves the good types indifferent to $(u_{g,h}^{FC}, u_{g,l}^{FC})$ but is still chosen in the presence of $(u_{g,h}^{FC}, u_{g,l}^{FC})$, must be less high-powered and would violate incentive compatibility, given that $(u_{b,h}^{FC}, u_{b,l}^{FC})$ is still the best contract for bad types by assumption.

remain active. If i 's unique best response is now withdrawal, set $\tilde{\sigma}_i^2(\tilde{s}^1) = W$ to arrive at the final $\tilde{\sigma}^2(\tilde{s}^1)$. It is a Nash equilibrium because $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, as constructed above, is the unique best response for firms $j \neq i, j \in \{3, 4\}$ if the deviator withdraws. It is in pure strategies by construction, and properties (i) and (ii) are empty due to $\Pi_i(\tilde{\sigma}^2(\tilde{s}^1)|\tilde{s}^1) = -\delta < 0$.

Assume now that $\delta = 0$. Construct $\tilde{\sigma}^2(\tilde{s}^1)$ from $\sigma^2(\tilde{s}^1)$ by first assuming that all $j \neq i, j \in \{1, 2\}$ play $\tilde{\sigma}_j^2(\tilde{s}^1) = NW$, which is always a best response for them, and initially keep all other players' strategies as in $\sigma^2(\tilde{s}^1)$. Even if this constitutes a change of strategy from $\sigma^2(\tilde{s}^1)$, the optimal behavior of non-deviators $j \neq i, j \in \{3, 4\}$ is clearly unaffected. If $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \neq \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for all $j \neq i, j \in \{3, 4\}$, indeed keep $\tilde{\sigma}_j^2(\tilde{s}^1) = \sigma_j^2(\tilde{s}^1)$ for them. Otherwise, if $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) = \Pi_j(W, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$ for at least one $j \neq i, j \in \{3, 4\}$, set $\tilde{\sigma}_j^2(\tilde{s}^1) = NW \forall j \neq i, j \in \{3, 4\}$. A similar argument as for the case $\delta > 0$ implies that they then give best responses against the profile constructed so far. If $\tilde{\sigma}_i^2(\tilde{s}^1) = NW$ is still a best response of the deviator, we have arrived at the desired equilibrium. Clearly, $\tilde{\sigma}^2(\tilde{s}^1)$ is in pure strategies, it has firms $j \neq i, j \in \{1, 2\}$ remaining active, and for any firm $j \neq i, j \in \{3, 4\}$ we can have $\tilde{\sigma}^2(\tilde{s}^1) = W$ only if $\Pi_j(NW, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1) < \Pi_j(W, \tilde{\sigma}_{-j}^2(\tilde{s}^1)|\tilde{s}^1)$, i.e. if there is no indifference. If, on the other hand, withdrawal is now the unique best-response of the deviator, setting $\tilde{\sigma}_i^2(\tilde{s}^1) = W$ yields the desired equilibrium, because if the deviator withdraws and $\delta = 0$, all firms $j \neq i, j \in \{1, \dots, 4\}$ are indifferent between withdrawing and remaining active, making the above constructed pure strategies best responses. Furthermore, the fact that $\Pi_i(\tilde{\sigma}^2(\tilde{s}^1)|\tilde{s}^1) = -\delta < 0$ implies that (i) and (ii) are empty. $\square$

For each $\tilde{s}^1 \in \tilde{S}^1$, replace the original Nash equilibrium $\sigma^2(\tilde{s}^1)$ with the pure-strategy equilibrium $\tilde{\sigma}^2(\tilde{s}^1)$.⁶⁸ In some of the corresponding subgames, using $\tilde{\sigma}^2(\tilde{s}^1)$ might already make the deviation unprofitable, i.e. $\Pi_i(\tilde{\sigma}^2(\tilde{s}^1)|\tilde{s}^1) \leq 0$. In fact, we show in the following that this is true in all $\Gamma^{FC}(\tilde{s}^1)$, $\tilde{s}^1 \in \tilde{S}^1$, if δ is sufficiently small. To prove this claim, we assume to the contrary that there are still profitable deviations. The stage 1b equilibria reached after these deviations do then satisfy the properties (i) and (ii) of Lemma 9. To save on notation, relabel the newly constructed stage 1b equilibria back to $\sigma^2(\tilde{s}^1)$, for all $\tilde{s}^1 \in \tilde{S}^1$, and, as before, let $\tilde{S}^1$ be the set of histories that still correspond to profitable unilateral deviations from s^1 by some firm $i \in \mathcal{J}$. For each $\tilde{s}^1 \in \tilde{S}^1$, denote by $\tilde{V}(\tilde{s}^1)$ the corresponding outcome in subgame $\Gamma^{FC}(\tilde{s}^1)$, i.e. the quadruple representing the two ex post types' choices among the available contracts at the end of stage 1b, and let $\tilde{d}(\tilde{s}^1)$ be the induced critical value for effort choice. $\tilde{V}(\tilde{s}^1)$ and $\tilde{d}(\tilde{s}^1)$ are well-defined because $\sigma^2(\tilde{s}^1)$ is in pure strategies.

Lemma 10. *There exists a value $\bar{\delta} > 0$ such that, if $0 \leq \delta < \bar{\delta}$, all outcomes $\tilde{V}(\tilde{s}^1)$, $\tilde{s}^1 \in \tilde{S}^1$, satisfy the constraints of FC.*

Proof. Consider any $\tilde{s}^1 \in \tilde{S}^1$. By definition of $\tilde{V}(\tilde{s}^1)$ as being the outcome in $\Gamma^{FC}(\tilde{s}^1)$ under $\sigma^2(\tilde{s}^1)$, it satisfies constraints (22) and (23). (25) must also be satisfied, because the offer (u_b, u_b) remains

⁶⁸If there are several equilibria that all satisfy the properties in Lemma 9 in a subgame $\Gamma^{FC}(\tilde{s}^1)$ for $\tilde{s}^1 \in \tilde{S}^1$, $\tilde{\sigma}^2(\tilde{s}^1)$ is just an arbitrary one of them.

active by construction of $\sigma^2(\tilde{s}^1)$.

Concerning (24), assume to the contrary that for some $\tilde{s}^1 \in \tilde{S}^1$, $\tilde{V}(\tilde{s}^1)$ violates (24), and let $\hat{S}^1 \subseteq \tilde{S}^1$ be the set of all such histories. As argued before, this implies losses for at least one active firm in $\Gamma^{FC}(\tilde{s}^1)$. Then, for each $\tilde{s}^1 \in \hat{S}^1$, let $\pi(\tilde{s}^1)$ be the (negative) profits of the active firm with the largest losses in $\Gamma^{FC}(\tilde{s}^1)$. We are going to show that there exists a value $\bar{\delta} > 0$ such that $\pi(\tilde{s}^1) \leq -\bar{\delta}$ for all $\tilde{s}^1 \in \hat{S}^1$, i.e. these losses are strictly bounded away from zero across all the histories $\tilde{s}^1 \in \hat{S}^1$.

Consider any $\tilde{s}^1 \in \hat{S}^1$. By assumption, $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$, and the non-deviators $j \neq i, j \in \{1, 2\}$ choose $\sigma_j^2(\tilde{s}^1) = NW$ and earn $\Pi_j(\sigma^2(\tilde{s}^1)|\tilde{s}^1) = 0$. Thus it must hold that V^{FC} satisfies (25) with slack, and for at least one $j \neq i, j \in \{3, 4\}$ it must be true that $\sigma_j^2(\tilde{s}^1) = NW$ and $\Pi_j(\sigma^2(\tilde{s}^1)|\tilde{s}^1) < 0$. As shown in the proof of Lemma 9, there are two cases in which this can happen. First, the deviator i 's best contract for good types in $\tilde{s}^1$ could be $(u_{g,h}^{FC}, u_{g,l}^{FC})$ and he does not offer a contract that is chosen by bad types in the presence of $(u_{b,h}^{FC}, u_{b,l}^{FC})$. Denote by $\hat{S}_1^1 \subset \hat{S}^1$ the set of deviation histories with this property. Second, the deviator i 's best contract for good types could be strictly preferred to $(u_{g,h}^{FC}, u_{g,l}^{FC})$ by good types. Let $\hat{S}_2^1 \subset \hat{S}^1$ be the set of histories in which this is the case. Hence $\hat{S}_1^1$ and $\hat{S}_2^1$ form a partition of $\hat{S}^1$.

Consider first a history $\tilde{s}^1 \in \hat{S}_1^1$. As we have shown in the proof of Lemma 9, the profits of an active non-deviator $j \neq i, j \in \{3, 4\}$ are then either β_1 or β_2 , depending on whether one or both of them are active non-deviators, with $\beta_1 < \beta_2 < 0$. Hence we know that $\pi(\tilde{s}^1) \leq \max\{\beta_1, \beta_2\} = \beta_2 < 0$ for all $\tilde{s}^1 \in \hat{S}_1^1$. Consider next a history $\tilde{s}^1 \in \hat{S}_2^1$ after which active non-deviators $j \neq i, j \in \{3, 4\}$ obtain only bad types. They earn $\pi_b(u_{b,h}^{FC}, u_{b,l}^{FC}) < 0$ with each unit mass of bad types agents that they obtain. Given that all bad types spread equally among at most three (and thus finitely many) firms, the losses $\pi(\tilde{s}^1)$ are strictly bounded away from zero across all $\tilde{s}^1 \in \hat{S}_2^1$ whenever the share of bad types $1 - G(\tilde{d}(\tilde{s}^1))$ is strictly bounded away from zero. But this is the case, because the good types' utility must be bounded above across profitable deviation histories $\tilde{s}^1 \in \hat{S}_2^1$, which implies that their share $G(\tilde{d}(\tilde{s}^1))$ must be strictly bounded away from one. Hence there exists a value $\beta_3 < 0$ such that $\pi(\tilde{s}^1) \leq \beta_3$ for all $\tilde{s}^1 \in \hat{S}_2^1$.

Putting the previous results together, we obtain that $\pi(\tilde{s}^1) \leq -\bar{\delta} := \max\{\beta_2, \beta_3\} < 0$ for all $\tilde{s}^1 \in \hat{S}^1$, i.e. whenever the outcome after a profitable deviation violates (24), a firm earns losses larger or equal to $\bar{\delta}$ in the corresponding stage 1b Nash equilibrium. But this is a contradiction to subgame perfection if $0 \leq \delta < \bar{\delta}$, because the firm would strictly prefer to withdraw, which implies our claim. $\square$

Hence if withdrawal costs are sufficiently small, the outcome after any profitable deviation must satisfy the constraints of FC. We next show that the outcome cannot be a solution to FC.

Lemma 11. *If $0 \leq \delta < \bar{\delta}$, it holds that $\tilde{V}(\tilde{s}^1) \notin \Omega^{FC}$ for all $\tilde{s}^1 \in \tilde{S}^1$.*

Proof. Assume to the contrary $\tilde{V}(\tilde{s}^1) \in \Omega^{FC}$ for some $\tilde{s}^1 \in \tilde{S}^1$. Given that solutions to FC can be Pareto ranked as discussed in Section 5.3, there are three possible cases. If $\tilde{V}(\tilde{s}^1)$ is considered worse

than V^{FC} by both ex post types, having $\tilde{V}(\tilde{s}^1)$ as outcome after the deviation requires $\sigma_j^2(\tilde{s}^1) = W \forall j \neq i, j \in \{3, 4\}$. Remaining active is, however, always a best-response for them in the presence of outcome $\tilde{V}(\tilde{s}^1)$, which implies that they actually choose $\sigma_j^2(\tilde{s}^1) = NW$ in $\sigma^2(\tilde{s}^1)$ by construction. If, second, $\tilde{V}(\tilde{s}^1) = V^{FC}$, and $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) > 0$, it must be true that V^{FC} satisfies (25) with slack, the deviator i has offered $(u_{g,h}^{FC}, u_{g,l}^{FC})$ but no contract chosen by bad types in the presence of $(u_{b,h}^{FC}, u_{b,l}^{FC})$, and $\sigma_j^2(\tilde{s}^1) = NW$ for at least one $j \neq i, j \in \{3, 4\}$. But then $\Pi_j(\sigma^2(\tilde{s}^1)|\tilde{s}^1) \leq -\bar{\delta}$, as shown in the proof of Lemma 10, which cannot occur in equilibrium if $\delta < \bar{\delta}$. Finally, if $\tilde{V}(\tilde{s}^1)$ is preferred to V^{FC} by all ex post types, it must have been offered by the deviator i , implying $\Pi_i(\sigma^2(\tilde{s}^1)|\tilde{s}^1) = 0$. $\square$

We thus know that, if $0 \leq \delta < \bar{\delta}$, after any profitable deviation history $\tilde{s}^1 \in \tilde{S}^1$ the outcome $\tilde{V}(\tilde{s}^1)$ in $\Gamma^{FC}(\tilde{s}^1)$ under $\sigma^2(\tilde{s}^1)$ must satisfy the constraints of FC but is not a solution to FC. Hence good types are strictly worse off in $\tilde{V}(\tilde{s}^1)$ than in V^{FC} , which requires that $\sigma_j^2(\tilde{s}^1) = W \forall j \neq i, j \in \{3, 4\}$. But if some firm $j \neq i, j \in \{3, 4\}$ remained active instead, it would earn non-negative profits $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \geq 0$. First, if it obtained bad types (in contract $(u_{b,h}^{FC}, u_{b,l}^{FC})$), the overall share of good types would be $G(\hat{d}^{FC})$, and all good types would choose $(u_{g,h}^{FC}, u_{g,l}^{FC})$. Even if firm j obtained all bad types, this ensures $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \geq 0$. Otherwise, firm j obtains only the good types (or no agents, if all agents decide to become bad types), also earning $\Pi_j(NW, \sigma_{-j}^2(\tilde{s}^1)|\tilde{s}^1) \geq 0$. Hence remaining active is a best response (even unique if $\delta > 0$), contradicting that $\sigma_j^2(\tilde{s}^1) = W$, by construction of $\sigma^2(\tilde{s}^1)$. This final contradiction shows that there cannot be profitable deviations if $0 \leq \delta < \bar{\delta}$.

Working Papers of the Socioeconomic Institute at the University of Zurich

The Working Papers of the Socioeconomic Institute can be downloaded from http://www soi.uzh.ch/research/wp_en.html

- 0814 Competitive Markets without Commitment, Nick Netzer, Florian Scheuer, November 2008, 65 p.
- 0813 Scope of Electricity Efficiency Improvement in Switzerland until 2035, Boris Krey, October 2008, 25 p.
- 0812 Efficient Electricity Portfolios for the United States and Switzerland: An Investor View, Boris Krey, Peter Zweifel, October 2008, 26 p.
- 0811 A welfare analysis of "junk" information and spam filters; Josef Falkinger, October 2008, 33 p.
- 0810 Why does the amount of income redistribution differ between United States and Europe? The Janus face of Switzerland; Sule Akkoyunlu, Ilja Neustadt, Peter Zweifel, September 2008, 32 p.
- 0809 Promoting Renewable Electricity Generation in Imperfect Markets: Price vs. Quantity Policies; Reinhard Madlener, Weiyu Gao, Ilja Neustadt, Peter Zweifel, July 2008, 34p.
- 0808 Is there a U-shaped Relation between Competition and Investment? Dario Sacco, July 2008, 26p.
- 0807 Competition and Innovation: An Experimental Investigation, May 2008, 20 p.
- 0806 All-Pay Auctions with Negative Prize Externalities: Theory and Experimental Evidence, May 2008, 31 p.
- 0805 Between Agora and Shopping Mall, Josef Falkinger, May 2008, 31 p.
- 0804 Provision of Public Goods in a Federalist Country: Tiebout Competition, Fiscal Equalization, and Incentives for Efficiency in Switzerland, Philippe Widmer, Peter Zweifel, April 2008, 22 p.
- 0803 Stochastic Expected Utility and Prospect Theory in a Horse Race: A Finite Mixture Approach, Adrian Bruhin, March 2008, 25 p.
- 0802 The effect of trade openness on optimal government size under endogenous firm entry, Sandra Hanslin, March 2008, 31 p.
- 0801 Managed Care Konzepte und Lösungsansätze – Ein internationaler Vergleich aus schweizerischer Sicht, Johannes Schoder, Peter Zweifel, February 2008, 23 p.
- 0719 Why Bayes Rules: A Note on Bayesian vs. Classical Inference in Regime Switching Models, Dennis Gärtner, December 2007, 8 p.
- 0718 Monoplistic Screening under Learning by Doing, Dennis Gärtner, December 2007, 29 p.
- 0717 An analysis of the Swiss vote on the use of genetically modified crops, Felix Schläpfer, November 2007, 23 p.
- 0716 The relation between competition and innovation – Why is it such a mess? Armin Schmutzler, November 2007, 26 p.
- 0715 Contingent Valuation: A New Perspective, Felix Schläpfer, November 2007, 32 p.
- 0714 Competition and Innovation: An Experimental Investigation, Dario Sacco, October 2007, 36p.
- 0713 Hedonic Adaptation to Living Standards and the Hidden Cost of Parental Income, Stefan Boes, Kevin Staub, Rainer Winkelmann, October 2007, 18p.
- 0712 Competitive Politics, Simplified Heuristics, and Preferences for Public Goods, Felix Schläpfer, Marcel Schmitt, Anna Roschewitz, September 2007, 40p.
- 0711 Self-Reinforcing Market Dominance, Daniel Halbheer, Ernst Fehr, Lorenz Goette, Armin Schmutzler, August 2007, 34p.

- 0710 The Role of Landscape Amenities in Regional Development: A Survey of Migration, Regional Economic and Hedonic Pricing Studies, Fabian Waltert, Felix Schläpfer, August 2007, 34p.
- 0709 Nonparametric Analysis of Treatment Effects in Ordered Response Models, Stefan Boes, July 2007, 42p.
- 0708 Rationality on the Rise: Why Relative Risk Aversion Increases with Stake Size, Helga Fehr-Duda, Adrian Bruhin, Thomas F. Epper, Renate Schubert, July 2007, 30p.
- 0707 I'm not fat, just too short for my weight – Family Child Care and Obesity in Germany, Philippe Mahler, May 2007, 27p.
- 0706 Does Globalization Create Superstars?, Hans Gersbach, Armin Schmutzler, April 2007, 23p.
- 0705 Risk and Rationality: Uncovering Heterogeneity in Probability Distortion, Adrian Bruhin, Helga Fehr-Duda, and Thomas F. Epper, July 2007, 29p.
- 0704 Count Data Models with Unobserved Heterogeneity: An Empirical Likelihood Approach, Stefan Boes, March 2007, 26p.
- 0703 Risk and Rationality: The Effect of Incidental Mood on Probability Weighting, Helga Fehr, Thomas Epper, Adrian Bruhin, Renate Schubert, February 2007, 27p.
- 0702 Happiness Functions with Preference Interdependence and Heterogeneity: The Case of Altruism within the Family, Adrian Bruhin, Rainer Winkelmann, February 2007, 20p.
- 0701 On the Geographic and Cultural Determinants of Bankruptcy, Stefan Buehler, Christian Kaiser, Franz Jaeger, June 2007, 35p.
- 0610 A Product-Market Theory of Industry-Specific Training, Hans Gersbach, Armin Schmutzler, November 2006, 28p.
- 0609 Entry in liberalized railway markets: The German experience, Rafael Lalive, Armin Schmutzler, April 2007, 20p.
- 0608 The Effects of Competition in Investment Games, Dario Sacco, Armin Schmutzler, April 2007, 22p.
- 0607 Merger Negotiations and Ex-Post Regret, Dennis Gärtner, Armin Schmutzler, September 2006, 28p.
- 0606 Foreign Direct Investment and R&D offshoring, Hans Gersbach, Armin Schmutzler, June 2006, 34p.
- 0605 The Effect of Income on Positive and Negative Subjective Well-Being, Stefan Boes, Rainer Winkelmann, May 2006, 23p.
- 0604 Correlated Risks: A Conflict of Interest Between Insurers and Consumers and Its Resolution, Patrick Eugster, Peter Zweifel, April 2006, 23p.
- 0603 The Apple Falls Increasingly Far: Parent-Child Correlation in Schooling and the Growth of Post-Secondary Education in Switzerland, Sandra Hanslin, Rainer Winkelmann, March 2006, 24p.
- 0602 Efficient Electricity Portfolios for Switzerland and the United States, Boris Krey, Peter Zweifel, February 2006, 25p.
- 0601 Ain't no puzzle anymore: Comparative statics and experimental economics, Armin Schmutzler, November 2008, 33p.
- 0514 Money Illusion Under Test, Stefan Boes, Markus Lipp, Rainer Winkelmann, November 2005, 7p.