

Dynamic Optimal Insurance and Lack of Commitment

Alexander K. Karaivanov and Fernando M. Martin
Simon Fraser University

June 25, 2007

Abstract

We analyze the role of commitment in a dynamic principal-agent model of optimal insurance with hidden effort and observable but non-contractible assets. We argue that the optimal contract under full commitment is time-inconsistent. Consequently, we solve for and analyze the optimal insurance contract when *both* the agent and the principal cannot commit (i.e., each can renege in each period) and contrast our results with the full commitment case studied by the existing literature. We find that the Markov-perfect contract under double-sided lack of commitment provides additional insurance relative to the self-insurance allocation and features a non-degenerate long-run asset and consumption distributions. Furthermore, the no-commitment contract differs significantly from the full commitment contract in the time profiles of consumption and savings, insurance, and welfare. We solve numerically for the optimal insurance contract in several environments characterized by different degrees of market imperfections. We find that the welfare loss due to lack of commitment is very high relative to the welfare costs of moral hazard or savings non-contractibility.

Keywords: optimal insurance, time-consistency, lack of commitment, moral hazard

JEL classification: D11, E21.

1 Introduction

We revisit the problem of optimal consumption smoothing and insurance provision in a stochastic multi-period principal-agent setting. The majority of the existing literature¹ assumes full commitment to a complete contract signed at time zero on behalf of both parties or at least the insurer. A typical result is that the optimal dynamic contract is characterized with a high degree of intratemporal consumption smoothing, but features a front-loaded consumption profile and extreme immiserization for the insured agent in the long run. The assumption of full commitment and the unpleasant immiserization result seem unattractive in many realistic situations. In addition, as we argue below, the full commitment optimal contract is time-inconsistent: that is, the insurer would like to change it at a later date, while the insured agent may like to deviate as well if her prescribed future discounted utility falls below the autarky value. Consequently, in this paper we analyze the problem of optimal dynamic insurance under lack of commitment while keeping in the background for comparison purposes the full commitment benchmark.

Specifically, we analyze an economy where a risk-averse agent is endowed with a stochastic output technology and can imperfectly self-insure through observable savings. Expected output is increasing in the agent's effort. We call this environment autarky. Alternatively, the agent can enter into an insurance contract with a risk-neutral principal. Due to the presence of asset accumulation, it is also clear that the problems of the agent and of the principal are both *dynamic* and the optimal contract is not just an outcome of a repeated game but is instead history dependent (Townsend, 1982) thus lack of commitment can prevent optimal intertemporal tie-ins.

We analyze and solve for the optimal insurance contract when both the agent and the principal cannot commit to the optimal time zero contract (i.e., they can renege each period) and contrast the results with the full commitment case. Our benchmark is a profit-maximizing principal, but we also look at the case of a benevolent planner. When analyzing contracts with lack of commitment, we focus on Markov-perfect equilibria which are the limit of finite horizon equilibria².

The autarky value of the agent plays a key role in our analysis since, in the optimal contract, the agent's participation constraint is binding. With full commitment the participation constraint binds at time zero only, whereas with no commitment, it binds in every period. We thus need to properly specify the value of autarky as a function of the agent's assets, instead of treating it as exogenous.

We first look at the effect of commitment on the consumption profile implied by the optimal contract. We start by analyzing the full information case, where the agent's effort is observable and contractible. In this case, it is always optimal for the risk neutral principal to provide full intratemporal insurance to the risk-averse agent. The first-best contract (under full commitment) features a (weakly) increasing marginal utility of consumption over time. Furthermore, the principal extracts all the agent's assets in the initial period. We find that the first best contract is still incentive feasible even if assets are non-contractible.

Under full commitment, the participation constraint of the agent is satisfied in the initial period only. Given that this constraint is binding and that the induced agent's consumption is front-loaded,

¹Green (1987), Thomas and Worrall (1990), Atkeson and Lucas (1992), Phelan and Townsend (1991), Phelan (1994) among many others.

²We thus rule out step-function equilibria as those obtained in Krusell and Smith (2003).

it follows that eventually the agent would be better off in autarky. Thus, the first-best contract is time-inconsistent, since the agent will want to walk away from the contract at some point. Hence we relax the commitment assumption and look at the optimal contract for the case where neither the principal nor the agent can commit to actions in future periods. In this case, we find that the participation constraint of the agent binds in every period. Since the principal is still providing full intratemporal insurance each period, the optimal contract features decreasing assets over time. However, in contrast to the commitment solution, once his assets are depleted, the agent still has to be given at least as much as he would earn in autarky; otherwise, he will walk away from the contract. Thus, in the long-run, the optimal contract with no commitment converges to a constant consumption level for the agent. Additionally, with lack of commitment, who controls asset accumulation matters unlike in the commitment case. If assets are non-contractible, the agent has an incentive to save away from zero assets as it improves his outside option and thus assets remain positive in the long-run.

We then analyze the optimal dynamic insurance contract under moral hazard due to hidden effort. Here, the principal cannot offer full intratemporal insurance due to the incentive problem. As in the full information case, the contract under commitment is time-inconsistent. However, we show that asset contractibility matters both under full and no commitment. When assets are contractible and there is full commitment, the optimal contract still features zero asset accumulation. In contrast, when assets are non-contractible, savings become an essential component of the optimal contract³. Without commitment, the participation constraint binds in every period, and thus the principal cannot front-load the transfers to the agent. However, we show that the principal still provides the agent with insurance above autarky. We also find that the full commitment and no-commitment contracts under moral hazard differ significantly in terms of the intertemporal profile of consumption, savings, intertemporal insurance and welfare. Finally, we briefly analyze the case of one-sided commitment, also known as “limited commitment”, where a contract is signed in the initial period, binding the principal’s actions, but allowing the agent to walk away at any point⁴.

The principal-agent contracting problem implies a Pareto frontier, on which the efficient utility realizations for the principal and the agent lie. Different institutional arrangements regarding the principal’s and agent’s ability to commit or regarding the contractibility of effort and savings affect the position of the frontier. Our approach allows us to identify the amount of insurance that can be optimally provided in various contractual environments, assess the value of commitment, and compare this value to the value of resolving moral hazard or asset contractibility problems. The profit-maximizing principal can be thought of as an insurance company, a revenue-maximizing politician, or even the mafia. The case of the benevolent planner is more related to macroeconomic issues like unemployment insurance, optimal taxation and social insurance.

Quantitatively, we find that, even without any commitment, the degree of insurance in the optimal contract far exceeds that in autarky, with consumption equivalence gains reaching almost 25% for the poorest agents under our benchmark parametrization. Alternatively, a profit maximiz-

³Thus, as shown in Fernandes and Phelan (2000) and Doepke and Townsend (2006), the correct state variable is not a scalar promised utility (as in a dynamic moral hazard problem without savings like in Phelan and Townsend, 1991 or Spear and Srivastava, 1987) but instead a *function* of assets. The intuition is the dynamic adverse selection problem introduced by the different incentives of agents who enter the next period with different asset levels (see Kocherlakota, 2004 for more details).

⁴For example, see Ligon, Thomas and Worrall (2002) who look at a model of mutual insurance among risk averse agents where each agent has the outside option to revert to autarky forever at any time. See also Phelan (1995).

ing principal could extract present value profits ranging up to 30% of average output in autarky. Changing the environment to allow resolution of the contractibility problems in assets and effort raises the welfare gains to 140% or profits more than twice average autarky output.

In the one-sided commitment case, the consumption equivalence compensation relative to autarky is at most 40% while under full commitment it is as high as 3,000%. The profits that a profit-maximizing principal can generate under commitment are as much as four times the average autarky output. We conclude that the value of commitment is potentially huge. The only exception is the case of profit-maximizing principal who can commit to the contract signed in the initial period while the agent cannot. There, the principal achieves the same profit as under the no-commitment case. The reason is that the agent's participation constraint is binding in both cases because committing to a promised utility for the agent higher than that in autarky, although possible for the principal, is suboptimal.

Related Literature

Optimal insurance schemes in dynamic moral hazard (unobserved action) settings were pioneered by Townsend (1982) in two-period and Rogerson (1985b) in multi-period models. These papers show that the optimal contract features "intertemporal tie-ins", i.e., it is not a simple repetition of the static contract as under the first best. The intuition is that the history of output realizations can be used to condition optimal insurance transfers to provide maximum expected welfare. The main implication to our model is that .

In contrast to ours, the above papers do not allow agents to save privately. This possibility was considered in Allen (1985) and Cole and Kocherlakota (2001), however for stochastic endowment economies. In Allen (1985), agents can privately save or borrow and he shows that no additional insurance over self-insurance can be provided. The intuition is that, regardless of his true (unobservable) history, an agent would always claim to have the history of income realizations that delivers the transfers with the highest present value. Cole and Kocherlakota (2001) strengthen Allen's result showing that, under some conditions⁵, it applies even if only private saving is allowed. Unlike Allen (1985) and Cole and Kocherlakota (2001), we obtain that additional insurance on top of what is achieved under self-insurance is always possible, even with lack of commitment. The reason is that we look at a moral hazard problem where the probability of different output realizations is endogenous. Thus, the policy of the planner always affects (through effort) the degree of uncertainty that the agent faces and welfare gains are possible unlike in the exogenous stochastic endowment setting⁶.

Until recently virtually all the literature on optimal dynamic contracting under private information has focused on the commitment case. For example, in all the above-mentioned papers, it is assumed that both the agent and the principal can commit to an optimal contract signed at time zero. In contrast, we are mostly interested in the case where the principal and the agent cannot commit. Our analysis thus relates to Ligon, Thomas and Worrall (2002) who look at limited commitment among a group of risk averse agents with stochastic endowments who mutually insure each other while each can at any moment revert to self-insurance through storage. They characterize the optimal contract and show that an improved storage technology can either raise

⁵A crucial condition is that the return on private savings is the same as the planner's return. Doepke and Townsend (2006) show that, if this condition is not satisfied, additional insurance is indeed possible.

⁶See also Abraham and Pavoni (2005).

or diminish welfare⁷. In contrast to their multi-agent setting where the only issue is insurance, we stick to the classic principal-agent formulation where a risk-neutral principal optimally insures a risk-averse agent facing an incentive problem due to moral hazard.

In our numerical methodology for solving the commitment case, we draw on methods introduced by Fernandes and Phelan (2000) and Doepke and Townsend (2006) and use promised utility as a function of assets together with the asset level of the agent as the state variables in our recursive formulation. However, in contrast with Doepke and Townsend’s linear programming approach, we show the validity and use the “first order approach” (Rogerson, 1985a, Werning, 2001, or Abraham and Pavoni, 2005). This allows us to use continuous variables (as opposed to grids) and increases the precision and computational speed. For the lack of commitment case, we use tools and methods developed in the literature on time-consistent government policy (Klein, Krusell and Ríos-Rull, 2005; Martin, 2006).

Our paper is also related to recent work by Acemoglu, Golosov and Tsyvinski (2006), who study a dynamic optimal contracting problem under lack of commitment in Mirrlees’ (1971) optimal taxation framework⁸. Our analysis of the no commitment case differs from theirs in several important dimensions. First, they allow for trigger strategies and thus focus on sustainable mechanisms, i.e., they look for outcomes that can be implemented through punishment strategies on the part of the agent. As is well-known, such equilibria are not re-negotiation proof and can only be implemented in infinite horizon economies. Instead, we focus on Markov equilibria which are the limit of finite-horizon equilibria. This allows us to analyze the fundamental equilibrium that arises under lack of commitment. Second, their problem is one of dynamic adverse selection (due to private information about agent’s productivity) so they need to ensure that a version of the revelation principle applies in order to solve for the optimal contract. In contrast, we avoid the difficulties with the potential inapplicability of the revelation principle under no commitment by differing from Acemoglu, Golosov and Tsyvinski in two directions: first, by focusing on a dynamic moral hazard problem (where hidden effort matters only within the period) and second, by assuming savings are observable but non-contractible. The observability of savings eliminates the need of eliciting truth-telling while their non-contractibility still preserves the dynamic nature of the problem and makes the Markov-perfect contract non-trivial.

2 The Model

We consider the problem of a long-lived risk-averse agent who derives utility from consumption, c and disutility from effort, e summarized by the concave function $u(c, e)$ satisfying standard Inada conditions. The agent discounts future utility by a factor $\beta \in (0, 1)$. He also produces output y , which he can consume or save. Output is stochastic and equals y^H with probability $\pi(e)$ and y^L with probability $1 - \pi(e)$ where $y^H > y^L > 0$ and $\pi(e)$ is increasing, concave and differentiable.

Given that the agent is risk averse and output is stochastic, the agent would like to smooth

⁷See also Kocherlakota (1996) who also analyzes (under perfect information) risk sharing between two risk-averse agents who cannot commit not to revert to autarky at any period. In stark contrast to our Markov approach and inefficiency results, he studies the set of subgame perfect equilibria and shows via a Folk theorem argument that, if agents are patient enough, there is no efficiency loss associated with the lack of commitment.

⁸See also Bisin and Rampini (2006) who look at a similar mechanism design problem with lack of commitment in a two period economy.

consumption over states of the world and time. The agent starts the period with some level of assets a , which can be carried over time via a linear saving technology that earns a fixed gross return, r . We assume that the agent cannot borrow, i.e., asset holding cannot be negative⁹. Thus, one possibility for the agent is to self-insure through savings which we refer to as *autarky*.

Clearly, self-insurance is insufficient to provide full consumption smoothing unless assets are infinitely large. Thus, if offered, the agent¹⁰ would demand additional insurance on top of what can be obtained through saving. Assume that there exists a risk neutral principal who can offer such insurance to the agent. The output realization is costlessly observable by the principal. The principal discounts his profits at rate R . We assume that

Assumption 1 $r \leq R \leq \frac{1}{\beta}$ and $r < \frac{1}{\beta}$.

The condition $r < \frac{1}{\beta}$ ensures that the autarky problem is well-defined, i.e., the agent does not accumulate an infinite amount of assets¹¹. Note also that we assume the principal to be weakly more patient than the agent. The discount rate of the principal can be interpreted in two basic ways. If $r = R$ then we can think of R as the rate at which the principal can move resources intertemporally, using the same technology as the agent. If $R = \frac{1}{\beta}$ then we can think of R as the rate of time preference for the planner, similar to that of the agent. For generality, we allow R to take any value between these two extremes.

As explained in the introduction, we analyze the optimal dynamic insurance contract that the risk-averse agent and the risk-neutral principal can enter into under various institutional environments: lack of commitment, non-contractible effort, non-contractible savings, as well as various combinations of these market imperfections. The insurance contract specifies history- and asset-contingent transfers (insurance premia and indemnities) to and from the agent. In all cases we assume that the principal observes the agent's asset level but the latter may or may not be contractible upon (i.e., the principal cannot necessarily force the agent to a particular asset level). This assumption makes the analysis tractable by eliminating the need for inducing truth-telling as well as avoiding known problems with the revelation principle under lack of commitment.

We start by solving the agent's problem under autarky which determines the outside option for the agent. No insurance contract that assigns time zero discounted utility lower than the autarky value is implementable. Furthermore, under lack of commitment by the agent, the principal must ensure that the latter obtains at least his autarky value in *each time period*. Thus, the autarky problem we study next serves a dual role: on the one hand it is a useful benchmark for the minimum amount of consumption smoothing the agent can achieve on his own and, on the other hand, it defines the agent's participation constraint in the optimal contract with the insurer.

⁹Alternatively, we could allow a negative lower bound on assets, e.g. the natural borrowing limit as in Aiyagari (1994). Our qualitative results are not affected by this assumption.

¹⁰Even though we talk about a single agent, our setting is more general as we solve the problem for any asset level, a . Thus, we can think of the principal as contracting separately with many agents with different initial savings.

¹¹See de Oliveira Sotomayor (1984) or Chamberlain and Wilson (2000).

2.1 The Agent's Problem In Autarky

The timing in each period is as follows. The agent first decides how much effort to put in, then output y is realized, and finally the agent decides how much to consume and save. Thus, depending on the state of the world (H or L), the agent faces the following budget constraint:

$$c^i + a^i = ra + y^i, \quad i = L, H$$

where a^i are assets carried over to the following period and c^i is current consumption, both conditional on the output realization.

Applying standard arguments, we can write the problem of the agent recursively as

$$\Omega(a) = \max_{a^H, a^L, e} \pi(e) (u(ra + y^H - a^H, e) + \beta\Omega(a^H)) + (1 - \pi(e)) (u(ra + y^L - a^L, e) + \beta\Omega(a^L))$$

subject to¹²

$$a^H, a^L \geq 0; \quad c^H, c^L > 0 \text{ and}$$

where $\Omega(a)$ is the agent's value function. It is immediate to notice that:

Lemma 1 *The value function in autarky, $\Omega(a)$, is increasing in assets.*

This result follows directly from the envelope condition of the autarky problem. Given that $r < \beta^{-1}$ by Assumption 1, the agent only saves to insure himself from consumption volatility. The higher his stock of savings, the more he can smooth out his consumption profile. Intuitively, an agent with a higher a can do everything an agent with lower a can but is in a better position to insure himself against a long enough sequence of low outputs, thus $\Omega_a(a) > 0$.

The agent's problem under autarky, which is essentially a well-known problem of self-insurance through savings, can be solved numerically using standard dynamic programming techniques. Since the constraint $a_L \geq 0$ may be binding for low asset levels, we solve the maximization problem directly instead of the system of first-order conditions. Computationally, we use an asset grid, but allow tomorrow's assets to take any value and interpolate between grid points using cubic splines.

As is standard in this type of models, the agent uses assets to insure against income risk. However, since markets are incomplete, self-insurance is not perfect, i.e., consumption in the high and the low income states differs. Assets are reduced if the agent is in the low income state and increased for some asset range if the agent is in the high income state. The condition $r < \frac{1}{\beta}$ ensures that the asset policy function in the high income states crosses the 45 degree line for some finite asset level and thus assets do not grow without bound¹³. This makes the autarky problem well-behaved and ensures convergence of the policy functions. Another important feature of the solution is that effort is decreasing in assets due to the agent being able to self-insure better at higher asset levels.

¹²We constrain effort on the open interval $(0, 1)$ to satisfy the full support assumption when we consider moral hazard later on and also to ensure $\pi(e)$ defines a proper probability measure.

¹³Note that this is exactly the opposite restriction of that used in Cole and Kocherlakota (2001). See their Proposition 4 for more details on this.

3 Optimal Dynamic Insurance Under Full Information

In this section we analyze and solve for the optimal dynamic insurance contract under full information, i.e., observable assets and full contractibility over the agent's effort. We start with the full information case since we view it as an intuitive benchmark we can use to assess directly the effect of lack of commitment on the optimal dynamic contract without introducing complications due to incentive-insurance trade-offs. As a baseline we look at a principal who maximizes expected net present value profits subject to a participation constraint for the agent. We discuss the case of a benevolent planner (competitive insurer) in the last section of the paper.

We start our analysis of the optimal insurance contract under full information with the case of full commitment. We show that, without loss of generality, the optimal contract features no asset accumulation by the agent, even if assets are non-contractible. We then study the case where both the principal and the agent cannot commit to an ex-ante contract at time zero. We demonstrate that with lack of commitment asset contractibility matters and the time profiles of consumption, savings and the degree of intertemporal insurance differ significantly from the commitment scenario.

As explained above, we look for contracts that give the agent at least his lifetime autarky value which represents the natural participation constraint for the agent. Depending on the degree of commitment we assume, this translates into either giving the agent at least the continuation value of autarky at time zero only (under full commitment) or in each period (under lack of commitment).

The timeline of the contracting problem is as follows. Given the insurance contract (transfers, savings and effort level), in the beginning of each period the agent applies the contracted effort level. Next, the output is realized and transfers take place. Finally, the agent consumes and decides on his asset level for tomorrow.

3.1 Full Commitment

We start our analysis with the first best where the agent's effort and savings are contractible and there is *full commitment* by both the insurer and the agent. This is a relatively well-known problem and we present it here solely as a benchmark. We then move on to the case of full commitment when the agent's assets are observable but not contractible.

3.1.1 The First Best

Let the history of output realizations up to time t be $s^t \equiv \{s_0, s_1, \dots, s_t\}$ where s_j is either L or H . The first best insurance arrangement is the solution to the following problem

$$\max_{\{\tau(s^t)\}, \{a(s^t)\}, \{e(s^{t-1})\}} \sum_{t, s^t} \frac{1}{R^t} \eta(s^t) (y(s_t) - \tau(s^t)) \quad (1)$$

subject to the participation constraint for the agent

$$\sum_{t, s^t} \beta^t \eta(s^t) u(c(s^t), e(s^{t-1})) \geq \Omega(a_0)$$

and a transversality condition

$$\lim_{t \rightarrow \infty} \beta^t a(s^t) = 0,$$

where $\tau(s^t)$ is the history dependent transfer from the planner to the agent, $\eta(s^t)$ is the probability of output history s^t , $c(s^t) = ra(s^{t-1}) + \tau(s^t) - a(s^t)$ is the agent's consumption, $a_0 \equiv a(s^{-1})$ is the agent's initial savings level, and $a(s^t) \geq 0$ for all t , s^t .

The first-order conditions with respect to¹⁴ τ_t^i , a_{t+1}^i and e_t are

$$-\frac{1}{R^t} + \beta^t \lambda u_c(c_t^i, e_t) = 0 \quad (2)$$

$$-u_c(c_t^i, e_t) + \beta r E u_c(c_{t+1}, e_{t+1}) \leq 0 \quad (3)$$

$$\sum_{i=\{L,H\}} \left\{ \frac{1}{R^t} \pi_e^i(e_t)(y^i - \tau^i) + \beta^t \lambda (\pi_e^i(e_t) u(c_t^i, e_t) + \pi^i(e_t) u_e(c_t^i, e_t)) \right\} = 0, \quad (4)$$

where λ is the multiplier on the participation constraint. The second inequality is strict if $a_{t+1}^i \geq 0$ binds.

From the first equation above we have

$$u_c(c_t^i, e_t) = \frac{1}{\lambda(R\beta)^t}, \quad (5)$$

which implies the agent's marginal utility of consumption is equalized across states for all t . Because effort, e_t is the same in both states, consumption is also equalized across states. The (standard) intuition is that, in the first best, the risk neutral principal provides full consumption smoothing over states of the world for the risk-averse agent as there are no incentive problems. This maximizes the surplus that can be generated.

Furthermore, notice that if the principal and the agent discount at the same rate, i.e., $R\beta = 1$, then the optimal marginal utility from consumption profile is constant over time as well, i.e., $u_c(c_t, e_t) = u_c(c_{t+1}, e_{t+1})$ for all t . In contrast, if $R < 1/\beta$, then marginal utility of consumption is rising over time. With separable utility (a sufficient but not necessary condition) this implies a falling consumption profile over time. The reason for this immiserization result is that the planner is more patient than the agent. Notice, however, that even if agent's consumption goes to zero (so under standard Inada conditions the agent's utility goes to minus infinity) the planner's profits are still bounded as they cannot exceed the value of extracting the high output level for sure and forever. The present utility value of the agent is also finite and equals the autarky value, $\Omega(a_0)$. Thus, a solution to the above problem always exists since autarky is always feasible.

Given the full insurance result, (5) implies

$$u_c(c_t, e_t) = \beta R u(c_{t+1}, e_{t+1}) \quad (6)$$

while (3) is equivalent to

$$u_c(c_t, e_t) \geq \beta r u_c(c_{t+1}, e_{t+1}). \quad (7)$$

¹⁴To simplify the notation, let $\tau_t^i \equiv \tau(s^t)$, $a_{t+1}^i \equiv a(s^t)$, $c_t^i \equiv c(s^t)$, and $e_t \equiv e(s^{t-1})$, when $s_t = i$, for $i = \{L, H\}$. Also define $\pi^i(e_t) \equiv \pi(e_t)$ if $i = H$ and $\pi^i(e_t) \equiv 1 - \pi(e_t)$ if $i = L$.

By Assumption 1 we have $r \leq R$. Notice that if $r = R$ then (6) implies (7), so both equations give exactly the same information about the time-path of consumption. Thus, assets are not necessary for the planner to implement the optimal allocation as this can be done through transfers only. This means that the initial assets of the agent can be extracted at any time (a form of Ricardian equivalence holds). On the other hand, if $r < R$, (6) implies (7) is satisfied with strict inequality and hence extracting the assets at time zero is strictly optimal when savings are contractible, since allowing the agent to carry assets at the low return r destroys surplus. Without loss of generality, we thus assume that assets are extracted in the initial period, i.e., $a_t^i = 0$ for all $t > 0$ in the first best contract. Interestingly, it turns out that the first best contract is still feasible (and hence optimal) even if assets non-contractible but still observable. These results are summarized in the following proposition:

Proposition 1 (a) *Without loss of generality, the optimal first best contract features no asset accumulation by the agent and a (weakly) increasing marginal utility of consumption profile.*
(b) *The first best contract in (a) is optimal even if the principal cannot control the agent's assets.*

Proof. (a) Shown above.

(b) Given that by Assumption 1 $r < 1/\beta$, the first-order conditions above imply that even if agent's savings are non-contractible, the first best allocation (c_t^*, e_t^*) is incentive compatible for the agent, i.e., he would not want to deviate from it. Indeed, suppose the agent wants to save one unit of resources. Note that the first order optimality condition for the agent's savings, a_{t+1} is identical to (7). Thus, the agent's loss of savings an extra unit in terms of today's utility is $u_c(c_t^*, e_t^*)$ while his gain tomorrow is $\beta r u_c(c_{t+1}^*, e_{t+1}^*)$. Clearly, from (6) and given $\beta r < 1$, the loss is (weakly) larger than the gain, thus the agent will never save if offered the first best allocation. But this means that the solution to the first best problem (which features no asset accumulation) is also the solution to the non-contractible savings problem where the agent could save but chooses not to. ■

The intuition for part (b) is that under full insurance the agent has no incentive to save. In autarky he chooses to save to self-insure but the need for that is clearly no longer present at the first best allocation. On the other hand, since $r \leq R$ the planner also has no incentive to leave the agent with any assets. This result simplifies the analysis of the contractible effort case, since we do not need to keep track of assets as a state variable. In contrast, as we show in the next section, in the case of lack of commitment asset accumulation plays a crucial role even when assets and effort are contractible.

3.1.2 Computing the Optimal Contract

Proposition 1 above and standard arguments from dynamic programming imply that the first best and the non-contractible savings optimal insurance contracts can be found by solving the two-stage recursive problem described below. As is typical in dynamic models of this type, the sufficient state variable is promised utility, w , that is, the discounted sum of future utility. In the first stage, the planner solves a static problem whereby he extracts the agent's initial assets, a_0 and promises

utility w' from the next period on, subject to the participation constraint for the agent:

$$\begin{aligned} \max_{\tau_0, e_0, w_1} \quad & \pi(e_0)y^H + (1 - \pi(e_0))y^L - \tau_0 + \frac{1}{R}\Pi(w_1) \\ \text{s.t.} \quad & u(ra_0 + \tau_0, e_0) + \beta w_1 = \Omega(a_0), \end{aligned}$$

The function $\Pi(\cdot)$ above is defined as the solution to the following (second-stage) dynamic programming problem

$$\begin{aligned} \Pi(w) \quad &= \max_{w', \tau, e} \pi(e)y^H + (1 - \pi(e))y^L - \tau + \frac{1}{R}\Pi(w') \\ \text{s.t.} \quad & u(\tau, e) + \beta w' = w. \end{aligned}$$

The solution to the above problems is the first best allocation $\{c_t^*, e_t^*\}$ implied by the sequence problem (14).

3.2 Lack of Commitment

In this section we keep the assumption of full information, i.e., that the agent's effort is observable and contractible but, unlike in the previous section, we assume that the principal cannot commit to future transfer policies. Thus, a contract needs to be signed in every period or, alternatively, the time-zero contract must be *time-consistent*. The agent also has limited commitment and can choose each period whether to participate or not in the contract, i.e., the optimal contract must give the agent at least the value of autarky each period.

More specifically we assume that, in the beginning of each period, both parties can decide to deviate from the previously specified arrangement as long as such deviations are profitable. For the agent, this involves comparing the continuation value of the insurance contract with his autarky value, $\Omega(a)$, which we derived in the previous section. We assume that such deviations are only possible *in the beginning* of each period before output is observed¹⁵. We thus think of the optimal contracting problem with lack of commitment as finding the best insurance contract such that at each period the agent obtains at least his autarky value and the principal has incentives to stick to a time-consistent profile of transfers.

We focus on Markov-perfect equilibria, which allow us to analyze contracts in which neither the agent nor the principal can commit to future choices. This includes ruling out the ability to commit to punishment strategies¹⁶. In particular, we do not allow the agent to threaten to revert to autarky forever, as such equilibria would not be renegotiation proof and thus not time consistent. In order to rule out the type of step function equilibria found by Krusell and Smith (2003), we look for Markov equilibria that are the limit of finite horizon equilibria.

The principal offers an insurance contract for the period based only on fundamentals and taking as given the policies that would be set by his future self. Given the agent's assets, a , the problem of

¹⁵Indeed, in our environment, if the principal and the agent could not commit even intra-temporally (i.e., within the period), the problem becomes trivial as only autarky can be supported. Specifically, if output is high the agent reverts to autarky while if output is low the principal reneges on the insurance transfer due.

¹⁶In this regard, we differ substantially from Acemoglu, Golosov and Tsyvinski (2006) who look at equilibria supported by punishments.

the principal today is to choose a transfer scheme $\{\tau^H, \tau^L\}$ conditional on output. The principal is taking into account that his future-self will follow some policy $\{\mathcal{T}^H(a), \mathcal{T}^L(a)\}$ that induces profits $\Pi(a)$ and that agent behavior is consistent with utility-maximizing.

The first-best contract analyzed in the previous section assumes that the principal and agent can commit to follow a contract signed at time zero. However, notice that this contract has the property that *both parties* would want to deviate from the specified arrangement if allowed to do so. In the case of the agent, the reason is that the participation constraint is only satisfied ex-ante. One implication of Proposition 1 is that at some point, the agent is left worse-off than under autarky. This means that the agent has no incentive to remain in the contract once he can do better by reverting to autarky. On the other hand, if the agent can be bound to stay in the contract forever, then the principal also has no incentive to follow the ex-ante agreed upon plan and would like to force the agent to work as hard as possible. Thus, the first-best contract is time-inconsistent. Note also that this result does not depend on whether assets are contractible or not.

Relaxing the commitment assumption implies that the participation constraint now needs to be satisfied in every period. Otherwise, the agent would not sign the contract and would revert to autarky for the period. Notice that we do not impose any restrictions on whether contracts are successfully signed or not in the future. As mentioned above, we do not allow neither the agent nor the principal to threaten to never again sign a contract if the current one is not to their satisfaction. Such threats are not credible since both parties would agree to sign a mutually advantageous contract at some future date, even if they “promise” not to do so today.

Formally, the policy implemented by the principal in the future implies a continuation value for the agent, $v(a)$. Given his future behavior, the principal today offers a contract which specifies transfers τ^H and τ^L , effort e , and future assets a^H and a^L that would leave the agent at least as well off as if he stays in autarky for the period. Thus, under no commitment, the participation constraint for the agent is

$$\pi(e) (u(ra + \tau^H - a^H, e) + \beta v(a^H)) + (1 - \pi(e)) (u(ra + \tau^L - a^L, e) + \beta v(a^L)) \geq \max_{\hat{a}^H, \hat{a}^L, \hat{e}} \pi(\hat{e}) (u(ra + y^H - \hat{a}^H, \hat{e}) + \beta v(\hat{a}^H)) + (1 - \pi(\hat{e})) (u(ra + y^L - \hat{a}^L, \hat{e}) + \beta v(\hat{a}^L)), \quad (8)$$

Note that the right hand side of the constraint depends only on a and $v(\cdot)$. Thus, we can write the participation constraint more compactly as

$$\pi(e) (u(ra + \tau^H - a^H, e) + \beta v(a^H)) + (1 - \pi(e)) (u(ra + \tau^L - a^L, e) + \beta v(a^L)) \geq \Phi(a; v). \quad (9)$$

The left hand side is just the agent’s value today, $v(a)$. Note that by construction $\Phi(a; v) \geq \Omega(a)$, i.e., going to autarky for one period weakly dominates staying in autarky forever.

3.2.1 Contractible Assets

Let us start with the case where the insurer can control the agent’s savings decisions. Given the future principal’s transfer policy $\{\mathcal{T}^H, \mathcal{T}^L\}$ inducing profits Π and agent behavior as summarized by v , the problem of today’s principal can be written recursively as

$$\Pi(a) = \max_{\tau^H, \tau^L, a^H, a^L, e} \pi(e) \left(y^H - \tau^H + \frac{\Pi(a^H)}{R} \right) + (1 - \pi(e)) \left(y^L - \tau^L + \frac{\Pi(a^L)}{R} \right)$$

subject to the participation constraint (9) (with Lagrange multiplier λ) and subject to $a^H, a^L \geq 0$. The following results allow us to simplify this problem.

Lemma 2 *The contract with full information and lack of commitment gives full intratemporal insurance to the agent, i.e., $\tau^H = \tau^L$.*

This follows immediately from the first-order conditions with respect to τ^H and τ^L . The standard intuition from the full commitment case still applies: the risk-neutral principal finds it optimal to fully insure the risk-averse agent as no incentive problems are present.

Proposition 2 *The participation constraint (9) is binding and $v(a) = \Omega(a)$.*

Proof. Note that the above Lemma implies

$$\lambda = \frac{1}{u_c},$$

which, given our assumptions on $u(c, e)$, implies $\lambda > 0$, i.e., the participation constraint is binding. Thus, $v(a) = \Phi(a; v) \geq \Omega(a)$. The participation constraint is binding no matter what value we put in its left hand side. The absolute minimum value the agent requires to accept the contract must provide present value utility weakly higher than staying in autarky forever. Given that the participation constraint is binding for any a and at any time, no $v(a) > \Omega(a)$ can be supported in equilibrium and thus $v(a) = \Omega(a)$. ■

The intuition for the proposition result is that the current principal leaves no surplus to the agent, regardless of how surplus is distributed in the future. In particular, if the principal tomorrow leaves the agent better-off than in autarky, today's principal will still extract all the surplus. In equilibrium, all principals act the same way and thus the agent is left just as well-off as in autarky.

The above results allow us to write the insurer's problem under lack of commitment as

$$\Pi(a) = \max_{\tau, a' \geq 0, e} \pi(e)y^H + (1 - \pi(e))y^L - \tau + \frac{\Pi(a')}{R}$$

subject to

$$u(ra + \tau - a', e) + \beta\Omega(a') - \Omega(a) = 0.$$

The first-order conditions for an interior solution are

$$\begin{aligned} -1 + \lambda u_c &= 0 \\ \frac{\Pi'_a}{R} + \lambda(-u_c + \beta\Omega'_a) &= 0 \\ \pi_e(y^H - y^L) + \lambda u_e &= 0. \end{aligned}$$

The envelope condition implies

$$\Pi_a = \lambda(ru_c - \Omega_a) = r - \frac{\Omega_a}{u_c}$$

Using the above, we can write the first-order conditions with respect to a' as

$$r - R + \Omega'_a \left(-\frac{1}{u'_c} + \frac{\beta R}{u_c} \right) = 0. \tag{10}$$

Proposition 3 *The optimal contract with lack of commitment and contractible effort and assets features an increasing marginal utility of consumption profile. In the long-run, assets are depleted, but consumption does not converge to zero in contrast to the full commitment case.*

Proof. From (10) we have

$$\frac{1}{u_c} = \frac{1}{\beta R u'_c} + \frac{R - r}{\beta R \Omega'_a}.$$

Given Assumption 1, we obtain

$$\frac{1}{u_c} > \frac{1}{\beta R u'_c} > \frac{1}{u'_c}$$

which (under separable utility) implies $c > c'$, i.e., agent's consumption decreases over time. What does c converge to? Applying Assumption 1, namely that $R \geq r$ and $\beta R \leq 1$ cannot hold at equality simultaneously, we have (since $\Omega_a > 0$ and $u_c > 0$):

$$r - R + \frac{\Omega_a}{u_c}(-1 + \beta R) < 0.$$

Thus, comparing with (10), we cannot have an interior steady state. In fact, the above expression indicates that the principal would like to decrease agent's assets as much as possible. Therefore, in the long-run, assets converge to their lower bound, zero and hence the long-run τ and e are characterized by

$$u(\tau, e) - (1 - \beta)\Omega(0) = 0 \quad \text{and} \quad \pi_e(y^H - y^L) + \frac{u_e}{u_c} = 0. \quad (11)$$

Our assumptions on u guarantee that $\tau > 0$ and so consumption is strictly positive in the long-run. ■

Compare the above result to the standard result from the full commitment model where consumption converges to zero in the long-run (full immiserization). With lack of commitment, we instead show that consumption converges to a strictly positive value in the long run. The reason for this difference is the participation constraint. Under lack of commitment, the participation constraint needs to be satisfied in each period since the agent can always walk away. Thus, even though the principal provides full intratemporal insurance and assets are consumed gradually, once assets are depleted, the principal can only retain the agent in the contract by offering him $\Omega(0)$ from then on, which implies offering positive consumption.

3.2.2 Non-Contractible Assets

An important question that comes up at this point is whether asset contractibility has an effect on the optimal contract under lack of commitment. Remember that in Proposition 1 we showed that, with full commitment, the first best contract is implementable even with non-contractible assets, as the agent has no incentive to save given the full insurance. We find that this result no longer holds under lack of commitment, namely the optimal insurance arrangements that can be achieved with and without asset contractibility are qualitatively different. The reason is that the agent wants to save more than the amount recommended by the principal in order to ensure higher expected transfers, since having more assets raises his outside option given by the autarky value.

We proceed to characterize the optimal contract under no commitment and non-contractible (but observable) assets. Because there is still no incentive problem and the agent is risk-averse, the optimal full information contract with lack of commitment still gives full insurance to the agent, even when assets are not contractible. Hence, given an announced transfer τ and contracted effort e , the agent's problem is

$$v(a) = \max_{a' \geq 0} u(ra + \tau - a', e) + \beta v(a').$$

The first-order condition is

$$-u_c + \beta v'_a \leq 0.$$

We use the first-order approach (no non-convexities here), so we take the above inequality as the incentive compatibility constraint for the principal. Note that, since the participation constraint is always binding, we can replace v'_a by Ω'_a . Moreover, since $a' \geq 0$, we can write the agent's optimality condition as

$$u_c - \beta \Omega'_a \geq 0.$$

The problem of the principal can be then written as

$$\Pi(a) = \max_{\tau, a' \geq 0, e} \pi(e)y^H + (1 - \pi(e))y^L - \tau + \frac{\Pi(a')}{R} \quad (12)$$

subject to

$$\begin{aligned} u(ra + \tau - a', e) + \beta \Omega(a') - \Omega(a) &= 0 \\ u_c - \beta \Omega'_a &\geq 0. \end{aligned}$$

Proposition 4 *In contrast to the full commitment case, the optimal contract with lack of commitment and non-contractible savings is different from that with contractible savings, i.e. it matters who controls asset accumulation.*

Proof. Let $\mu \geq 0$ be the Lagrange multiplier of the constraint $u_c - \beta \Omega'_a \geq 0$ so we have $\mu(u_c - \beta \Omega'_a) = 0$. Suppose the contractible and non-contractible savings contracts are the same. Then, $\mu = 0$ and $u_c - \beta \Omega'_a \geq 0$ (the ICC cannot be strictly binding at the contractible assets allocation). Recall the first-order condition with respect to a' in that case:

$$\frac{\Pi'_a}{R} + \lambda(-u_c + \beta \Omega'_a) = 0.$$

Since $\lambda = 1/u_c > 0$, then $u_c - \beta \Omega'_a \geq 0$ implies $\Pi'_a \geq 0$, i.e., profits have to be weakly *increasing* for all a .

From the steady state condition (11)

$$u(\tau, e) = (1 - \beta)\Omega(0).$$

we see that τ is such that the agent is indifferent between getting τ for sure in every period and getting stochastic draws of y^H and y^L , given that he starts with zero assets. Since u is strictly concave, the agent is willing to pay an insurance premium. By Jensen's inequality we have that τ (in this case, the certainty equivalent consumption level) is lower than expected consumption

in autarky. Thus, the profits of the principal must be strictly positive for $a = 0$. Given our assumptions on preferences, the demand for insurance, and therefore the insurer's profits, are decreasing in assets. In particular, profits go to zero as assets go to infinity since the agent is able to perfectly self-insure. Thus, Π must be *decreasing* for some high enough asset levels, a contradiction. ■

Let us further characterize the optimal dynamic insurance contract in this case. With Lagrange multipliers λ and μ for the first and second constraints, respectively, the first-order conditions are

$$\begin{aligned}\frac{\Pi'_a}{R} + \lambda(-u_c + \beta\Omega'_a) - \mu(u_{cc} + \beta\Omega'_{aa}) &= 0 \\ -1 + \lambda u_c + \mu u_{cc} &= 0 \\ \pi_e(y^H - y^L) + \lambda u_e + \mu u_{ce} &= 0.\end{aligned}$$

From Proposition 4 we know that $\mu > 0$, otherwise the problem is equivalent to the contractible savings case. Thus, given $u_c - \beta\Omega'_a = 0$, the first-order conditions above simplifies to

$$\frac{\Pi'_a}{R} - \mu(u_{cc} + \beta\Omega'_{aa}) = 0. \quad (13)$$

To continue further, assume that utility is separable, i.e. $u_{ce} = 0$. We have

$$\begin{aligned}\lambda &= -\frac{\pi_e(y^H - y^L)}{u_e} \\ \mu &= \left(1 + \frac{u_c\pi_e(y^H - y^L)}{u_e}\right) \frac{1}{u_{cc}}.\end{aligned}$$

In addition, from the envelope condition:

$$\Pi_a = \frac{\Omega_a\pi_e(y^H - y^L)}{u_e} + r$$

thus we can re-write (13) as

$$u_c(y^H - y^L) \left(\frac{\pi'_e}{\beta R u'_e} - \frac{\pi_e}{u_e} \right) + \frac{r}{R} - 1 - \beta \frac{\Omega'_{aa}}{u_{cc}} \left(1 + \frac{u_c\pi_e(y^H - y^L)}{u_e} \right) = 0.$$

The solution for the optimal $\{\tau, a', e\}$ is then characterized by the above equation and the two constraints of the problem (12).

How about the long-run? The equations characterizing a steady state are

$$\begin{aligned}\frac{u_c\pi_e(y^H - y^L)}{u_e} \left(\frac{1}{\beta R} - 1 \right) + \left(\frac{r}{R} - 1 \right) - \beta \frac{\Omega_{aa}}{u_{cc}} \left(1 + \frac{u_c\pi_e(y^H - y^L)}{u_e} \right) &= 0 \\ u((r-1)a + \tau, e) - (1-\beta)\Omega(a) &= 0 \\ u_c - \beta\Omega_a &= 0.\end{aligned}$$

Proposition 5 *In the long-run, the allocation implied by the optimal insurance contract under lack of commitment and non-contractible savings is interior and generically features positive assets.*

Proof. Focus on the first of the equations characterizing the steady state. Since $\beta R \leq 1$, the first term is zero if $R\beta = 1$ and negative otherwise. The term $\frac{r}{R} - 1$ may be zero or negative. However, by Assumption 1 both terms cannot be zero simultaneously. Thus, the first two terms are jointly negative.

To get an interior solution, we need the third term of the equation to be positive. Given that $\mu > 0$, we know that

$$1 + \frac{u_c \pi_e (y^H - y^L)}{u_e} < 0.$$

Therefore, since Ω and u are strictly concave, the last term in the equation is indeed always positive and the solution in the long-run is thus interior. Zero assets could be a solution, but this is a non-generic outcome. ■

4 Optimal Contracting Under Moral Hazard

In this section we consider the problem of dynamic optimal insurance when the agent's effort is unobservable so a moral hazard problem ensues. It is well known that in this case the optimal contract cannot feature full insurance since insurance has to be traded off against incentive provision. The moral hazard problem introduces an extra constraint (incentive compatibility) in the planner's problem. It is easy to verify (see Karaivanov, 2006) that in our setting with two output levels and concave π the monotone likelihood ratio property (MLRP) and the convexity of the distribution function condition hold and thus the "first-order approach" (Rogerson, 1985a) is valid. Thus, everywhere below we replace the incentive constraint with the first-order condition with respect to effort.

4.1 Lack of Commitment

Consider the problem of a profit-maximizing principal with no commitment and non-contractible assets and effort. As before, given that the principal in the future will follow a transfer policy $\{\mathcal{T}^H, \mathcal{T}^L\}$ inducing profits Π and agent behavior as summarized by the value function v , the current principal's problem is

$$\Pi(a) = \max_{\tau^H, \tau^L, e, a^H, a^L} \pi(e) \left(y^H - \tau^H + \frac{1}{r} \Pi(a^H) \right) + (1 - \pi(e)) \left(y^L - \tau^L + \frac{1}{r} \Pi(a^L) \right)$$

subject to the agent maximizing utility

$$\{a^H, a^L, e\} = \arg \max_{a^H, a^L, e} \pi(e) (u(ra + \tau^H - a^H, e) + \beta v(a^H)) + (1 - \pi(e)) (u(ra + \tau^L - a^L, e) + \beta v(a^L)),$$

the agent's participation constraint

$$\begin{aligned} & \pi(e) (u(ra + \tau^H - a^H, e) + \beta v(a^H)) + (1 - \pi(e)) (u(ra + \tau^L - a^L, e) + \beta v(a^L)) \geq \\ & \max_{\hat{a}^H, \hat{a}^L, \hat{e}} \pi(\hat{e}) (u(ra + y^H - \hat{a}^H, \hat{e}) + \beta v(\hat{a}^H)) + (1 - \pi(\hat{e})) (u(ra + y^L - \hat{a}^L, \hat{e}) + \beta v(\hat{a}^L)), \end{aligned}$$

and the non-negativity constraints

$$a^H, a^L, \hat{a}^H, \hat{a}^L \geq 0.$$

As before, we look for a Markov-perfect equilibrium, i.e., a set of functions $\{\mathcal{T}^H, \mathcal{T}^L, \Pi, \mathcal{A}^H, \mathcal{A}^L, \mathcal{E}, v\}$ that solves the above problem. Note that solving the problem involves finding a fixed point in the unknown functions $\Pi(a)$ and $v(a)$.

The principal's problem as written above is cumbersome. In particular, the participation constraint involves a maximization problem since we need to check for one-period reversal to autarky. Fortunately, as in the case with full information, the participation constraint is always binding. The intuition is the same whether there is moral hazard or not. A profit-maximizing principal has no incentive to leave the agent with any of the surplus of the contract. Moreover, because of the lack of commitment, dynamic considerations do not matter either: if a principal tomorrow leaves some agent better-off than in autarky, today's principal will still extract all the surplus. Thus, in equilibrium the principal leaves the agent exactly as well-off as in autarky. This means that the value to the agent of signing the contract, $v(a)$, is equal to his value under autarky, $\Omega(a)$ which simplifies the problem considerably, since $\Omega(a)$ is relatively easy to compute.

As explained above, to further simplify the problem, we use the first-order approach. Thus, we assume that the equilibrium is differentiable (almost everywhere). We already know that the problem of the agent in autarky is well behaved and that its value function is differentiable. However, we need to be careful with the first-order condition with respect to a^L , as the non-negativity constraint on a^L may be binding (see more on this below).

Given that the profit-maximizing principal in the future will follow a policy $\{\mathcal{T}^H, \mathcal{T}^L\}$ that induces profits Π , the simplified optimal contracting problem is

$$\Pi(a) = \max_{\tau^H, \tau^L, e, a^H, a^L} \pi(e) \left(y^H - \tau^H + \frac{1}{r} \Pi(a^H) \right) + (1 - \pi(e)) \left(y^L - \tau^L + \frac{1}{r} \Pi(a^L) \right)$$

subject to

$$\begin{aligned} -u_c(c^H, e) + \beta \Omega_a(a^H) &= 0 \\ -u_c(c^L, e) + \beta \Omega_a(a^L) &\leq 0 \\ \pi(e) (u(c^H, e) - u(c^L, e) + \beta(\Omega(a^H) - \Omega(a^L))) + \pi(e) u_e(c^H, e) + (1 - \pi(e)) u_e(c^L, e) &= 0 \\ \pi(e) (u(c^H, e) + \beta \Omega(a^H)) + (1 - \pi(e)) (u(c^L, e) + \beta \Omega(a^L)) &= \Omega(a) \end{aligned}$$

and non-negativity constraint¹⁷ $a^L \geq 0$, where $c^i = ra + \tau^i - a^i$ for $i = L, H$.

The solution to this problem is a set of functions $\{\mathcal{T}^H, \mathcal{T}^L, \Pi, \mathcal{A}^H, \mathcal{A}^L, \mathcal{E}\}$. Unfortunately, because of the non-linear constraints, analytical results are hard to obtain and the optimal contract with moral hazard and lack of commitment needs to be solved numerically. We explain the details in the next section.

One important question that our approach allows us to answer quantitatively is how much additional profits can the principal make by having more control over the agent's behavior, i.e.,

¹⁷We omit $a^H \geq 0$, since it is not binding. Note that the second constraint above is binding if $a^L > 0$ and is negative if $a^L = 0$.

allowing savings, effort or both to be contractible. Note that solving these cases is not difficult once we have solved the general problem written above, as it simply amounts to de-activating the appropriate incentive-compatibility constraints. In particular, to compute the contractible assets regime, we take out the first two constraints given by the Euler equations of the agent's problem (i.e., the first-order conditions with respect to assets); for the contractible effort environment, we take out the first-order condition with respect to effort, and for the contractible assets and effort regime we de-activate all three incentive compatibility constraints.

4.2 Full Commitment

We now consider the case when the two parties can fully commit to an insurance contract at time zero. This will allow us to contrast our results under no commitment obtained above with the better studied case of full commitment.

4.2.1 Sequential Formulation

Assume first that the agent's savings are contractible. Proceeding as in section 3 and using the first order approach, the optimal insurance contract solves

$$\max_{\{\tau(s^t)\}, \{a(s^t)\}, \{e(s^{t-1})\}} \sum_{t, s^t} \frac{1}{R^t} \eta(s^t) (y(s^t) - \tau(s^t)) \quad (14)$$

subject to

$$\begin{aligned} \sum_{t, s^t} \beta^t \eta(s^t) u(c(s^t), e(s^{t-1})) &\geq \Omega(a_0) \\ \sum_{i=L, H} \left\{ \pi_e^i(e_t) (u(c_t^i, e_t) + \beta \omega^i(s^t)) + \pi^i(e_t) u_e(c_t^i, e_t) \right\} &= 0 \quad \forall t, \end{aligned}$$

where

$$\omega^i(s^t) \equiv \sum_{j=1}^{\infty} \sum_{s^{t+j}|s^t} \beta^{j-1} \frac{\eta(s^{t+j})}{\eta(s^{t-1}) \pi^i(e_t)} u(c(s^{t+j}|s^t), e(s^{t+j}|s^t)),$$

and subject to the transversality condition, and the non-negativity constraints on assets. The first constraint is the (ex-ante) participation constraint and the second constraint is the (per-period) incentive constraint with respect to agent's effort, where $\omega^i(s^t)$ denotes future expected discounted utility given history s^t .

We proceed to characterize the optimal dynamic insurance contract in this case. In particular, we are interested whether, as in the no moral hazard case, a contract extracting all assets in the initial period is still optimal which, if true, would simplify the solution methodology significantly. The answer is given by the following proposition.

Proposition 6 (a) *Without loss of generality the optimal contract under full commitment, moral hazard, and contractible savings features zero asset accumulation by the agent.*

(b) *If $u_{ec} \leq 0$ (leisure and consumption are weak complements) and $r = R$ then the optimal*

insurance contract from (a) is not incentive compatible with the agent's individual savings decision, i.e. the agent is always savings-constrained at the allocation implied by the moral hazard contract and thus the optimal contract under full commitment, moral hazard, and non-contractible assets generically features asset accumulation by the agent.

Proof. (see Appendix) ■

Note several important implications of the above proposition.

(1) If we maintain the assumption that $u_{ec} \leq 0$, then the optimal allocation in part (a) where all assets are taken away in the initial stage can be compatible with non-contractible savings only if R is larger than r and large enough (or, r being small enough), so that it is possible to have $u_c \leq \beta R E u'_c$ and $u_c > \beta r E u'_c$ hold simultaneously. Intuitively, if the “public” return on assets, R , is large enough, then the agent can be (optimally) made to voluntarily relinquish his ability to save.

(2) Allowing consumption and leisure to be substitutes, $u_{ec} > 0$ may, if strong enough, reverse the above results since the wedge between u_c and $R\beta E u'_c$ narrows and possibly changes sign.

(3) The result in part (b) shows that, unlike all other full commitment contractual environments we studied, asset accumulation is likely to be an essential part of the optimal contract under non-contractible savings and moral hazard. Importantly, this implies that the relevant state variable capturing history in the recursive formulation of the sequence problem above *must be* $w(a)$, i.e., promised utility as a *function of assets*, as discussed in the introduction (see Fernandes and Phelan, 2000 or Doepke and Townsend, 2006). Indeed, suppose the correct state variable for promises was instead a scalar w independent of a . Then, given that w' would not depend on a' , the agent will have no incentive to save in the current period and the contractible assets allocation would be implementable - a contradiction.

(4) Suppose $u_{ec} = 0$, i.e., preferences are separable in consumption and leisure. Then, (17) implies that, if $\pi_e > 0$ (which is true for the high state), the right hand side of (17) evaluated at the first best consumption is smaller than $1/\lambda$ (its value under the first best), hence c^H under moral hazard is higher than the first best level and by symmetry c^L will be lower. That is, only partial insurance obtains as usual.

Finally, take a somewhat closer look at the problem of the principal under moral hazard and non-contractible assets. This problem is the same as above, with the two additional (per-period) constraints coming from the agent's optimal savings condition

$$u_c(c_t^i, e_t) - \beta r E u(c_{t+1}, e_{t+1}) \geq 0,$$

for all $i = L, H$, and all t . The following Lemma characterizes more sharply the optimal contract in this case.

Lemma 3 *The consumption profile under full commitment, moral hazard and non-contractible assets is such that consumption in each state at time t is lower than its corresponding value under moral hazard and contractible assets.*

Proof. (see Appendix) ■

4.2.2 Recursive Formulation

To solve for the optimal contract, we once again write the problem recursively using the techniques proposed by Spear and Srivastava (1987) by keeping track of promised utility as the state variable reflecting any possible history dependence. The presence of asset accumulation in our problem creates an additional intertemporal link and necessitates the use of different promises for different asset levels (see Fernandes and Phelan, 2000; Doepke and Townsend, 2006), i.e., the states of our dynamic programming problem are a and some promised utility *function*, $w(a)$ that belongs to a set of feasible functions, $\mathcal{W}$, which has to be endogenously defined and iterated on together with the value function in the numerical implementation. Notice that these promise functions are not present in the no-commitment problem as they are exactly the mathematical representation of the idea of commitment.

Given the above, in the recursive formulation of the principal's problem, the principal chooses transfers and promised utility functions $w^H(a)$ and $w^L(a)$ which belong to a set of feasible functions $\mathcal{W}$. Thus, for any $a \geq 0$ and any $w(\cdot) \in \mathcal{W}$, the problem of the principal is

$$\begin{aligned} \Pi^c(a, w(a)) = & \max_{\tau^i, a^i, w^i(\cdot), e} \pi(e) \left(y^H - \tau^H + \frac{1}{R} \Pi^c(a^H, w^H(a^H)) \right) \\ & + (1 - \pi(e)) \left(y^L - \tau^L + \frac{1}{R} \Pi^c(a^L, w^L(a^L)) \right) \end{aligned} \quad (15)$$

subject to the agent optimizing

$$\begin{aligned} \{a^H, a^L, e\} = & \arg \max_{a^H, a^L, e} \pi(e) (u(ra + \tau^H - a^H, e) + \beta w^H(a^H)) \\ & + (1 - \pi(e)) (u(ra + \tau^L - a^L, e) + \beta w^L(a^L)), \end{aligned}$$

promise keeping

$$\pi(e) (u(ra + \tau^H - a^H, e) + \beta w^H(a^H)) + (1 - \pi(e)) (u(ra + \tau^L - a^L, e) + \beta w^L(a^L)) = w(a),$$

and

$$a^H, a^L \geq 0, w^H(\cdot), w^L(\cdot) \in \mathcal{W}$$

Solving the above full commitment problem is much harder than its corresponding no-commitment problem from the previous section since now we need to iterate on the set $\mathcal{W}$, which contains the feasible promise functions $w(a)$ using the methods from Abreu, Pierce and Stacchetti (1990). In practice, one starts with a large enough set which contains the fixed point of the feasibility correspondence defined by the constraints of the above problem and iterates on the resulting set together with iterating on the value function Π^c . Natural lower and upper bounds to the set $\mathcal{W}$ can be defined by bounding consumption and/or assets —see the next section for more details.

Compare the optimal contract under full commitment solving the above problem to that of the no-commitment case from section 4.1. The first difference is that now the principal has to satisfy the participation constraint only in the initial period. Thus, the optimal way to maximize profits is to promise the agent a lower utility on average for tomorrow, potentially differentiating between high and low output. As the agent becomes utility-poorer, he works more, which allows the principal to extract a larger surplus. These incentives are present for any given promise function $w(a)$.

Now, look at the incentives of the agent, given the principal’s incentives to lower promised utility over time. Since the agent’s value function is increasing in his assets a , the agent has an incentive to accumulate as much assets as possible which is the only way for him to counteract the principal. In other words, we have two conflicting forces in the above problem: the principal wants to lower w as much as possible by front-loading agent’s consumption while at the same time the agent wants to “hedge” himself against this by accumulating savings and thus guaranteeing himself higher future promises. Because of this, the numerical solution of the problem (15) appears unbounded in the optimal policies¹⁸ for w^i and a^i (the principal wants to lower w as much as possible while the agent wants to raise a as much as possible) and hence the value function iteration does not converge unless we set an exogenous lower bound on promises and an upper bound on assets.

To get more intuition for this numerical difficulty, look back at the problem with commitment where both assets and effort are contractible. In this case, as seen in section 3, the principal can extract the agent’s assets in the first period without loss of generality and so we can solve the simpler problem without assets avoiding the computational problems described above.

5 Numerical Analysis

In this section we discuss the methods one can use to solve for the optimal insurance contracts under no or full commitment numerically. We also compute a full set of results for a benchmark parametrization and assess the value of commitment versus asset or effort contractibility. A comprehensive set of robustness checks is available at the authors’ website.

5.1 Parametrization

For computational purposes we will restrict ourselves to the following functional forms

$$\begin{aligned} u(c, e) &= \frac{(\alpha c^\rho + (1 - \alpha)(1 - e)^\rho)^{\frac{1-\sigma}{\rho}} - 1}{1 - \sigma} \\ \pi(e) &= e^\nu, \quad e \in (0, 1) \end{aligned}$$

Preferences have the flexible CES form, where $\rho = 1$ corresponds to consumption and leisure being perfect substitutes, $\rho = 0$ corresponds to Cobb-Douglas preferences, and $\rho = -\infty$ corresponds to perfect complements. Table 5.1 shows the benchmark parameter values we use in the simulations.

Table 1: Benchmark parameter values

α	ρ	σ	β	ν	r	y^H	y^L
0.70	0.00	1.50	0.93	0.50	1.04	0.30	0.10

¹⁸Our numerical computation s where we allowed the principal to choose from a discrete set of promise functions $w(\cdot)$ parallel to each other, showed that the principal wants to choose very low promises very quickly. Since he has to satisfy the promise keeping constraint, this implies a sharp increase in assets (due to promised utility being increasing in assets) and so the problem “explodes” in both directions.

5.2 Lack of commitment

5.2.1 Moral Hazard and Non-contractible Assets

To compute the solution to the problem with no commitment, moral hazard and non-contractible assets, we use an asset grid, but allow variables, including tomorrow's assets, to take any value. To interpolate between grid points we use cubic splines. The numerical methodology we use follows that described in Martin (2006) and basically involves finding a fixed-point for the profit function $\Pi(a)$: start with a guess for $\Pi(a)$; solve the principal's problem, which outputs a new profit function; update and continue iterating until there is convergence.

Figures 1 to 5 show the computed solution at our benchmark parameters. In summary, the principal is able to provide insurance in addition to the self-insurance available to the agent under autarky. The degree of additional insurance is decreasing in the agent's assets and so are the profits of the principal. Figure 1 shows that $\tau^H < y^H$ and $\tau^L > y^L$ for all asset levels. If realized output is low, then the principal runs a deficit and provides the agent with higher income than in autarky; if realized output is high, then the principal makes a profit by providing the agent with less income than in autarky. As shown in Figure 3, this implies that compared to autarky, consumption is lower in the high income state and higher in the low income state. Since the participation constraint is binding, the agent is as well-off as in autarky, which allows the principal to make a profit; effectively, the agent is paying a premium for the added insurance. Because of the extra insurance, assets in the high income state grow slower than in autarky. The extra income in the low state also allows the agent to save a bit more in that case — see Figure 2.

Figure 4 compares effort under autarky to that in the insurance contract with lack of commitment. In both cases effort is decreasing in the agent's assets, but effort is higher in autarky. This is an immediate consequence of the moral hazard problem. As we argued above, profits are decreasing in assets (see Figure 5). The intuition is that the agent needs less insurance when he is wealthier.

5.2.2 The Value of Asset and Effort Contractibility

How much additional profits can the principal make by having more control over the agent's behavior? We consider the following cases: (1) full information with contractible assets (FIC), (2) full information with non-contractible assets (FIN), (3) moral hazard with contractible assets (MHC), and (4) moral hazard with non-contractible assets (MHN). This last case is the one we just analyzed above and we will use it as benchmark. From section 3, we also know some of the properties of the first two cases. In particular, we know that under full information, there is full intratemporal insurance. Furthermore, in the FIC case assets converge to zero in the long run, whereas we get an interior steady state for FIN.

Figure 6 shows the resulting expected discounted profits for the four regimes listed above. Relative to the benchmark case (MHN), giving the principal control of asset accumulation (MHC) increases his profits, but not significantly. On the other hand, if the principal could observe the agent's effort (FIN) his profits increase dramatically (they triple in our parametrization). Adding further control of asset accumulation (FIC) now has a big effect. The resulting profits in the case where both assets and effort are contractible (FIC) are in the order of 5 to 6 times larger compared

to the baseline with all the constraints active (MHN).

5.3 Commitment

Given the unboundedness issues with the commitment problem under moral hazard and non-contractible assets explained in section 4.2, in order to be able to perform comparisons with the lack of commitment case, instead of solving for the optimal profit value of the principal as defined above, we solve a simplified problem that provides a *lower bound* on the profits that can be achieved under commitment. In particular, assume that the lowest possible promised utility that a principal can support is the value of promising the agent consumption equal to the low output level, y^L and zero effort forever. Notice that, given the moral hazard problem, if consumption is set to be constant at y^L , then $e = 0$ is the only incentive compatible effort level. Let the value of the agent from this allocation (allowing him to adjust assets optimally) be $\Phi(a)$.

Now assume that the principal follows the following policy: starting at some initial state $(a, \Omega(a))$, set $w^H(a) = w^L(a) = \Phi(a)$ and set the rest of the controls such that the agent's ex-ante utility is equal to that in autarky, $\Omega(a)$ so that the participation constraint is satisfied ex-ante. Then, after “jumping” to the promise function $\Phi(a)$ the principal keeps promising this value forever and chooses transfers and effort and asset level recommendations to satisfy the resulting promise-keeping constraint. Clearly, this policy/contract is suboptimal since it does not allow rewarding high output with a higher promise or gradual reduction in promises between $\Omega(a)$ and $\Phi(a)$. Still, it turns out, the profits realized by following this suboptimal policy swamp anything that could be achieved under no commitment.

Figure 7 compares expected net present value profits for selected lack of commitment cases with the profits for the full commitment case with moral hazard and non-contractible savings obtained by following the suboptimal policy described above. As we can see, with full commitment, the principal can extract much more profit than under no commitment, even if in the latter case he was able to control the agent's savings and effort decision. Specifically, with commitment, the principal makes about 2.5 times as much than under the no-commitment optimal arrangement with contractible a and e . This example shows that it can be much more valuable for a principal to secure a commitment technology than controlling all the agent's decisions. Remember also that the commitment line on the figure is actually a lower bound of what the true profits under commitment would be, so the true value of commitment is even higher.

5.4 Long-Run Properties

Table 2 shows some long-run properties of the different cases we have considered so far. The table does not make any welfare statements. It is just a description of how an economy would look like in the long-run. Using the solutions obtained from the previous sections we compute some aggregate statistics after running the model for a 100 periods for 100,000 agents. For all cases considered, the economy converges quite faster than the 100 periods run.

The case with lack of commitment and moral hazard distinguishes itself from the others, since the asset distribution in the limit is non-degenerate. Figure 8 shows the Lorenz curves for the cases of autarky, MHN and MHC. Clearly, when the agent faces uncertainty about output or transfers,

Table 2: Simulation of 100,000 agents for 100 periods. Sample averages of 100th period.

	Autarky	MHN	MHC	FIN	FIC	Commit.
τ	0.225	0.213	0.212	0.215	0.215	0.140
e	0.394	0.331	0.333	0.372	0.410	0.472
a/y	0.874	0.672	0.424	0.509	0.000	0.443
c/y	1.036	1.016	1.004	0.985	0.941	0.609
π/y	0.000	0.011	0.014	0.035	0.059	0.410

the more insurance is provided, the more unequal the asset distribution becomes. This result is in stark contrast with the cases of full information and commitment. With full information (FIN and FIC), the principal provides full intratemporal insurance; thus, all agents hold the same level of assets in the limit. With commitment, there is full immiserization; thus assets converge to zero for all agents.

6 Extensions

6.1 Benevolent Planner

In this section we expand our results by analyzing the case of a benevolent planner who maximizes the expected discounted utility of the agent. In the contractual relationship, this corresponds to giving all the bargaining power to the agent as opposed to the principal having all the bargaining power as before. In terms of the resulting optimal contracts, much of the intuition from sections 3 and 4 carries through, so we focus on the differences.

Compared to its profit maximizing counterpart, the benevolent planner faces a different constraint: zero expected present value profits. Note that the participation constraint is no longer binding, since all the surplus of the contract goes to the agent and autarky is always feasible. For computational reasons, we do not allow for the planner to accumulate its own assets¹⁹. As before, we first analyze the case of no commitment and then compare the results with the full commitment benchmark.

6.1.1 Lack of Commitment

When the principal is unable to commit to his future actions, the zero expected present value profit constraint simplifies to a zero expected profits per period constraint. The reason is simple: since every planner in the future has zero expected present value profits, the current planner must have zero expected profits for the period too. Thus, in equilibrium, all planners have zero expected profits per period. Note that this argument would not hold if we allowed the planner to accumulate

¹⁹Basically, there is an incentive for the planner to acquire the assets of the agent and accumulate them himself. Even if both discount at the same rate, the planner would have full control over his own assets, whereas it needs to satisfy incentive-compatibility constraints for the agent's assets.

assets.

Given that the benevolent planner in the future will follow some policy $\{\mathcal{T}^H, \mathcal{T}^L\}$ that induces agent behavior as summarized by $\mathcal{V}$, the problem of the principal is

$$\mathcal{V}(a) = \max_{\tau^H, \tau^L} \pi(e) (u(ra + \tau^H - a^H, e) + \beta\mathcal{V}(a^H)) + (1 - \pi(e)) (u(ra + \tau^L - a^L, e) + \beta\mathcal{V}(a^L))$$

subject to the agent optimizing

$$\begin{aligned} -u_c(c^H, e) + \beta\mathcal{V}_a(a^H) &= 0 \\ -u_c(c^L, e) + \beta\mathcal{V}_a(a^L) &\leq 0 \\ \pi_e(e) (u(c^H, e) - u(c^L, e) + \beta(\mathcal{V}(a^H) - \mathcal{V}(a^L))) + \pi(e) u_e(c^H, e) + (1 - \pi(e))u_e(c^L, e) &= 0, \end{aligned}$$

zero expected period profits

$$\pi(e)(y^H - \tau^H) + (1 - \pi(e))(y^L - \tau^L) = 0,$$

and the non-negativity constraint $a^L \geq 0$.

As before, we define a Markov-perfect equilibrium to be the set of functions $\{\mathcal{T}^H, \mathcal{T}^L, \mathcal{A}^H, \mathcal{A}^L, \mathcal{E}, \mathcal{V}\}$ that solves the above problem. To solve for the equilibrium we need to find a fixed-point in $\mathcal{V}(a)$.

We can solve the above problem numerically, using the methods described in section 5. Let us compare the contracts under the benevolent planner and the profit-maximizing principal. In terms of differences, expected transfers and consumption in both income states are higher under the benevolent planner. This is not surprising since all the surplus of the contract now goes to the agent. The higher consumption also translates in lower asset accumulation.

There are also similarities between the profit maximizing and benevolent planner environments without commitment. Most salient is the fact that the effort function is remarkably similar. This implies that the “size of the pie”, as measured by the level of expected output, is similar for any level of assets. However, given that asset accumulation is lower under the benevolent planner, expected output increases faster over time. Therefore, expected present value output is actually higher under the benevolent planner than under the profit-maximizing principal.

The benevolent planner case allows us to ask a natural welfare question, namely how much more current consumption would the agent need in autarky to be indifferent between staying in autarky and signing-up with the benevolent planner? This involves finding, for every a , the fraction ζ of current consumption under autarky such that the agent is indifferent between the two environments, i.e., solving

$$\pi(e) (u((1 + \zeta)c^H, e) + \beta\Omega(a^H)) + (1 - \pi(e)) (u((1 + \zeta)c^L, e) + \beta\Omega(a^L)) = \mathcal{V}(a).$$

The gains for the agent from “signing-up” with the benevolent planner are as high as 24% of consumption when his assets are zero and about 4% of consumption when his assets are at their maximum value of 2. These welfare gains appear quite substantial, especially for the poorer agents.

As in case of the profit-maximizing principal, we also analyze the implications of relaxing the endogenous market incompleteness due to the non-contractibility of assets and/or effort. Figure 9 shows the consumption equivalence compensation, ζ , for the environments in which it is possible for

the planner to control assets, effort, or both. It is clear that under a benevolent planner, virtually all gains come from eliminating the informational friction, i.e., making effort contractible. When the planner controls both assets and effort, the gains from the optimal insurance contract relative to autarky can be as much as 140% of consumption. Even when assets are high the gains are considerable: about 27% of consumption at $a = 2$.

6.1.2 Full Commitment

We use a similar solution strategy to that employed for the profit-maximizing principal case in order to assess the welfare gains possible under commitment. Again, due to the computational difficulties explained in the previous section, we study a suboptimal policy rule for the principal that bounds the true gains from commitment from below. Specifically, we once again use the lower bound on promises defined in section 4.2., namely the value, $\Phi(a)$ of consuming y^L forever. We first compute the value of planner's profits from promising $\Phi(a)$ forever to the agent, independently of the realized output. We then solve for the initial promise level (which clearly has to be higher than the autarky value) which makes the present value of profits for the planner at $t = 0$ equal to zero. Figure 10 shows the consumption equivalence compensation for the full commitment benevolent planner case. Clearly, the gains in this case are far larger than anything that can be achieved under lack of commitment. Even at our lower bound, welfare gains are as much as 3,000% of autarky consumption.

6.2 One-Sided Commitment

This final section briefly discusses the case of one-sided commitment which lies in between the full and no commitment environments studied above. Specifically, assume that the principal still specifies all his future actions in the initial period and is able to commit. However, unlike in the full commitment case, the agent can walk away from the contract at any time period. Thus, in the one-sided commitment scenario, we restrict $w(a) \geq \Omega(a)$ for every period, by which we mean that we set the lower bound of the set $\mathcal{W}$ to be the autarky value function, $\Omega(a)$. The initial period corresponds to the state $w(a) = \Omega(a)$.

We argued that under full commitment, the profit-maximizing principal would like to lower agent's promises over time. Because of the incentive to promise less than today and the constraint of $w(a) \geq \Omega(a)$ in every period, we expect to have $w^H(a) = w^L(a) = \Omega(a)$. Thus, the solution of the one-sided commitment problem should be identical to the no commitment solution. Our numerical results confirm this intuition.

In the benevolent planner case, to illustrate the properties of the optimal insurance contract we solve for the initial promise for each asset level, $w(a)$ which, choosing optimally the transfers and asset and effort recommendations and setting $w^H(a) = w^L(a) = \Omega(a)$ from the first period on and forever, gives a zero present value of profits for the principal. That is, we first solve for the optimal contract when the agent is always promised autarky independent of the output realization and then, knowing the profits of the planner (some positive number) for each a we pick an initial promise, $w(a)$ to make the present value of profits zero. Figure 11 shows the resulting consumption equivalence compensation. The welfare gains for the agent are larger than under lack of commitment with no asset or effort contractibility (MHN), but smaller than when the planner can control both asset

accumulation and effort (FIN).

7 Conclusions

...

References

- [1] Abraham, A. and N. Pavoni (2005), "Efficient Allocations with Moral Hazard and Hidden Borrowing and Lending", mimeo, University of Rochester and UCL.
- [2] Abreu, D., D. Pearce and E. Stacchetti (1990), "Toward a Theory of Discounted Repeated Games with Imperfect Monitoring", *Econometrica* 58(5), pp. 1041-63.
- [3] Acemoglu, D., M. Golosov and A. Tsyvinski (2006), "Markets Versus Governments: Political Economy of Mechanisms", working paper, MIT.
- [4] Aiyagari, S.R., (1994), "Uninsured Idiosyncratic Risk and Aggregate Saving," *Quarterly Journal of Economics*, 109(3), pp. 659-84.
- [5] Allen, F., (1985), "Repeated principal-agent relationships with lending and borrowing," *Economics Letters*, 17(1-2), pp. 27-31.
- [6] Atkeson, A. and R. Lucas, (1992), "On efficient distribution with private information", *Review of Economic Studies*, 59, pp. 427-453
- [7] Bisin, A. and A. Rampini (2006), "Markets as beneficial constraints on the government," *Journal of Public Economics*, 90(4-5), pp. 601-629.
- [8] Chamberlain, G. and C. Wilson (2000), "Optimal Intertemporal Consumption under Uncertainty," *Review of Economic Dynamics*, 3(3), pp. 365-395.
- [9] Cole, H. and N. Kocherlakota (2001), "Efficient Allocations with Hidden Income and Hidden Storage," *Review of Economic Studies*, 68(3), pp. 523-42.
- [10] de Oliveira Sotomayor, Marilda A. (1984), "On Income Fluctuations and Capital Gains," *Journal of Economic Theory*, 32(1), pp. 14-35.
- [11] Doepke, M. and R. Townsend (2006), "Dynamic mechanism design with hidden income and hidden actions," *Journal of Economic Theory*, 126(1), pp. 235-285.
- [12] Fernandes, A. and C. Phelan (2000), "A Recursive Formulation for Repeated Agency with History Dependence", *Journal of Economic Theory*, 91, pp. 223-247.
- [13] Green, E. (1987), "Lending and Smoothing of Uninsurable Income", in E.C. Prescott and N. Wallace (eds), *Contractual Arrangements for Intertemporal Trade*, pp. 3-25. University of Minnesota Press.
- [14] Karaivanov, A. (2006), "Financial Constraints and Occupational Choice in Thai Villages", working paper, Simon Fraser University.
- [15] Klein, P., P. Krusell and J.-V. Ríos-Rull (2006), "Time-Consistent Public Expenditures", mimeo.
- [16] Kocherlakota, N., (1996), "Implications of Efficient Risk Sharing without Commitment", *Review of Economic Studies*, 63(4), pp. 595-609.

- [17] Kocherlakota, N., (2004), "Figuring out the Impact of Hidden Savings on Optimal Unemployment Insurance," *Review of Economic Dynamics*, 7(3), pp. 541-554.
- [18] Krusell P. and A. Smith (2003), "Consumption-Savings Decisions with Quasi-Geometric Discounting," *Econometrica*, 71(1), 365-375.
- [19] Ligon, E., J. Thomas and T. Worrall, (2002) "Informal Insurance Arrangements with Limited Commitment: Theory and Evidence from Village Economies," *Review of Economic Studies*, 69(1), pp. 209-44.
- [20] Martin, F. (2006), "Optimal Taxation without Commitment," mimeo, Simon Fraser University.
- [21] Mirrlees, J. (1971), "An Exploration in the Theory of Optimum Income Taxation," *Review of Economic Studies*, 38(114), pp. 175-208.
- [22] Phelan, C. (1995), "Repeated Moral Hazard and One-Sided Commitment," *Journal of Economic Theory*, 66(2), pp. 488-506.
- [23] Phelan, C. and R. Townsend (1991), "Computing Multi-Period, Information-Constrained Equilibria", *Review of Economic Studies*, 58, pp. 853-881.
- [24] Rogerson, W. (1985a), "The First Order Approach to Principal-Agent Problems", *Econometrica* 53, pp. 1357-68.
- [25] Rogerson, W. (1985b) "Repeated Moral Hazard," *Econometrica* 53, pp. 69-76.
- [26] Spear, S. and S. Srivastava, (1987), "On Repeated Moral Hazard with Discounting", *Review of Economic Studies*, 54(4), pp. 599-617.
- [27] Thomas, J. and T. Worrall, (1990), "Income fluctuation and asymmetric information: An example of a repeated principal-agent problem," *Journal of Economic Theory*, 51(2), pp. 367-390.
- [28] Townsend, R. (1982), "Optimal Multi-period Contracts and the Gain from Enduring Relationships Under Private Information", *Journal of Political Economy* 61, pp. 166-86.
- [29] Werning, I. (2001), "Repeated Moral-Hazard with Unmonitored Wealth: A Recursive First-Order Approach", manuscript, University of Chicago.

Appendix

Proof of Proposition 6

(a) Clearly this must be the case when the planner controls savings (the pure moral hazard case) because we assume $r \leq R$ and it is sub-optimal for the planner to leave assets with the agent. Thus, the only possible reason to leave assets with the agent in this case is if this somehow relaxes the agent's incentive constraint with respect to effort. However, the latter constraint is affected by the agent's consumption and effort profiles not by assets per se. Suppose we assume that the optimal contract features some asset levels held by the agent $\tilde{a}_{t+1}^i$ and some transfers $\tilde{\tau}_t^i$. Then we can always re-define $\tau_t^i = \tilde{\tau}_t^i + r\tilde{a}_t - \tilde{a}_{t+1}^i$ and $a_{t+1}^i = 0$ for all t , and implement exactly the same consumption profile for the agent, thus keeping his incentives intact. Moreover, the new transfers yield weakly higher profits for the planner given that asset extraction occurs immediately.

(b) We will show that the agent would like to save away from the allocation from (a). While the proof below can be done using the sequence problem directly, it is easier to work with the equivalent recursive formulation. The recursive problem solved by the moral hazard allocation *with no asset accumulation* from (a) is²⁰

$$\Pi(w) = \max_{\tau^i, w^i, e} \sum_{i=L, H} \pi^i(e)(y^i - \tau^i + \frac{1}{R}\Pi(w^i))$$

subject to

$$\begin{aligned} \sum_{i=L, H} \pi^i(e)(u(c^i, e) + \beta w^i) &= w \\ \sum_{i=L, H} \left\{ \pi_e^i(e)(u(c^i, e) + \beta w^i) + \pi^i(e)u_e(c^i, e) \right\} &= 0, \end{aligned}$$

With Lagrange multipliers λ and μ , the first-order condition with respect to w^i is

$$\pi^i(e) \left(\frac{1}{R}\Pi_w(w^i) + \lambda\beta + \mu\beta \frac{\pi_e^i(e)}{\pi^i(e)} \right) = 0,$$

which simplifies to

$$\Pi_w(w^i) + \lambda\beta R + \mu\beta R \frac{\pi_e^i(e)}{\pi^i(e)} = 0.$$

From the envelope condition, $\Pi_w(w^i) = -\lambda^i$, where λ^i is the multiplier value tomorrow, given output y^i today. Thus,

$$\lambda^i = \beta R \left(\lambda + \mu \frac{\pi_e^i(e)}{\pi^i(e)} \right).$$

Multiply both sides by π^i , add up and get

$$E(\lambda') = \lambda\beta R, \tag{16}$$

²⁰The allocation in the initial period must be solved as shown after Proposition 1 by defining the first-stage static problem.

since $\pi_e^L = -\pi_e^H$ (the probabilities add up to 1). The above implies that if $R\beta = 1$ then λ is a martingale which is a common result in this type of models (see Atkeson and Lucas, 1991 for the case where the income process is exogenous).

Now take the first-order condition with respect to τ^i . We have

$$-\pi^i(e) + \lambda\pi^i(e)u_c(c^i, e) + \mu\pi_e^i(e)u_c(c^i, e) + \mu\pi^i(e)u_{ec}(c^i, e) = 0,$$

which reduces to

$$u_c(c^i, e) = \frac{1 - \mu u_{ec}^i}{\lambda + \mu \frac{\pi_e^i}{\pi^i}}. \quad (17)$$

If $u_{ec} \leq 0$ and replacing λ using (16) we get

$$\lambda^i \geq \frac{\beta R}{u_c(c^i, e)} \quad (18)$$

Then, using (18) and Jensen's inequality we have:

$$\frac{\beta R}{\lambda} = \frac{(\beta R)^2}{E(\lambda')} \leq \frac{\beta R}{E\left(\frac{1}{u_c(c, e)}\right)} \leq \beta R E u_c(c, e) \quad (19)$$

Look at the agent's incentives to save at the allocation implied above. As in Proposition 1, the utility loss from saving one unit of consumption is $u_c(c, e)$ and the gain is $\beta r E u_c(c', e')$. The expression in (18) and the inequality (19) taken one period forward, imply that u_c today is smaller than $\beta R E u'_c$ at the moral hazard constrained allocation. Thus, if $r = R$, as in our benchmark specification, (19) implies that the gain from saving a unit of consumption away from the moral hazard constrained allocation is always larger than the cost, i.e., the agent would always like to save away. In other words, the optimal moral hazard constrained allocation can be implemented only if the agent's ability to save is controlled by the principal. Notice that this result is a generalization in a multi-period setting of the similar result in Rogerson (1985b).

Proof of Lemma 3

With γ_t^i as the Lagrange multiplier of each of the constraints above, the first-order conditions of the moral hazard problem with non-contractible assets imply

$$u_c^i = \frac{\frac{1}{R^t \beta^t} - \mu u_{ec}^i - \gamma_t^i \frac{u_{cc}}{\pi^i}}{\lambda + \mu_t \frac{\pi_e^i}{\pi^i}}.$$

Given that $u_{cc} < 0$, the above implies that, evaluated at the moral hazard allocation with contractible assets, the right hand side is larger than u_c^i which means that c_t^i under non-contractible assets must be *lower* than its corresponding value under contractible assets.

Figure 1: Profit-maximizing principal — No commitment: Transfers

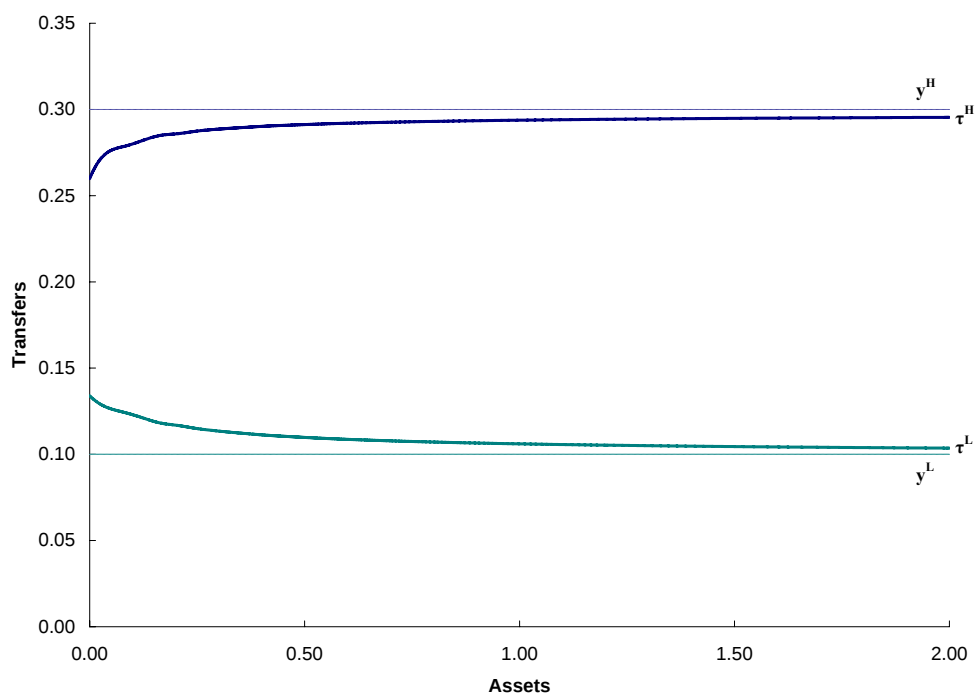


Figure 2: Profit-maximizing principal — No commitment: Assets

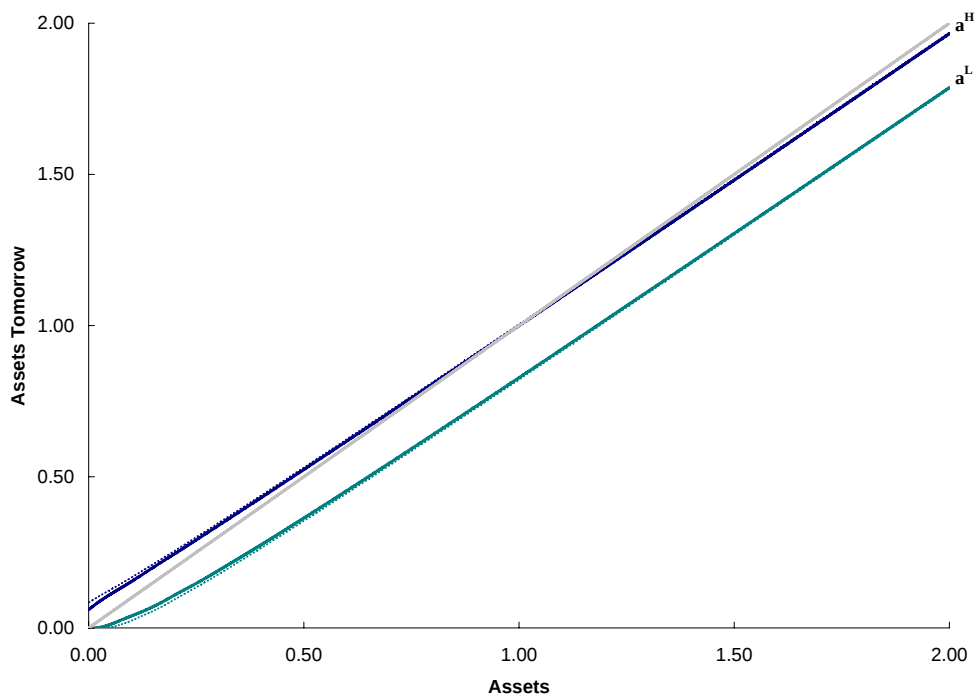


Figure 3: Profit-maximizing principal — No commitment: Consumption

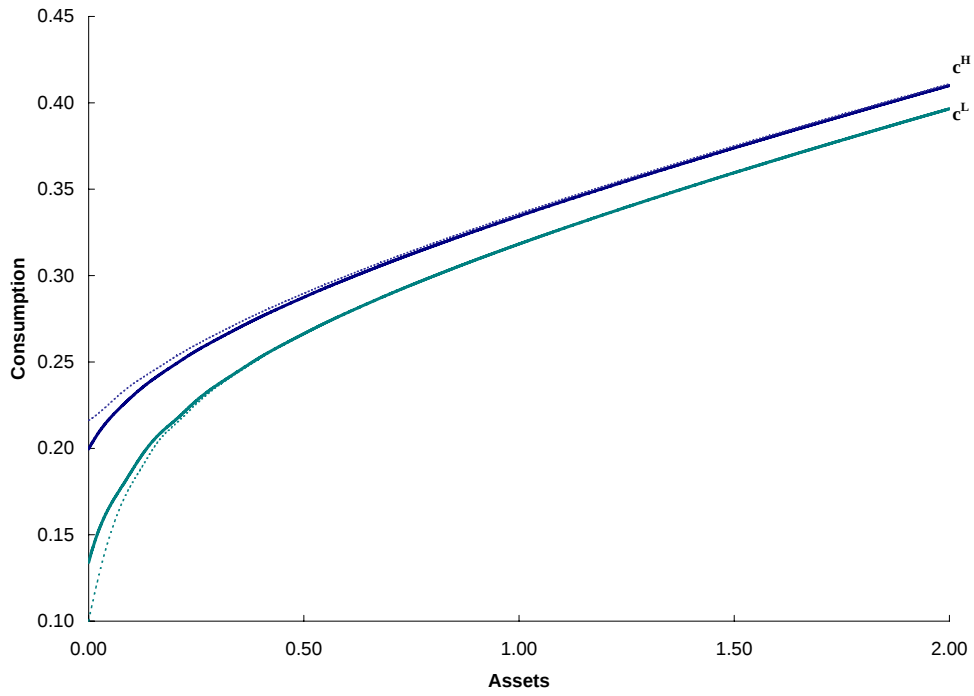


Figure 4: Profit-maximizing principal — No commitment: Effort

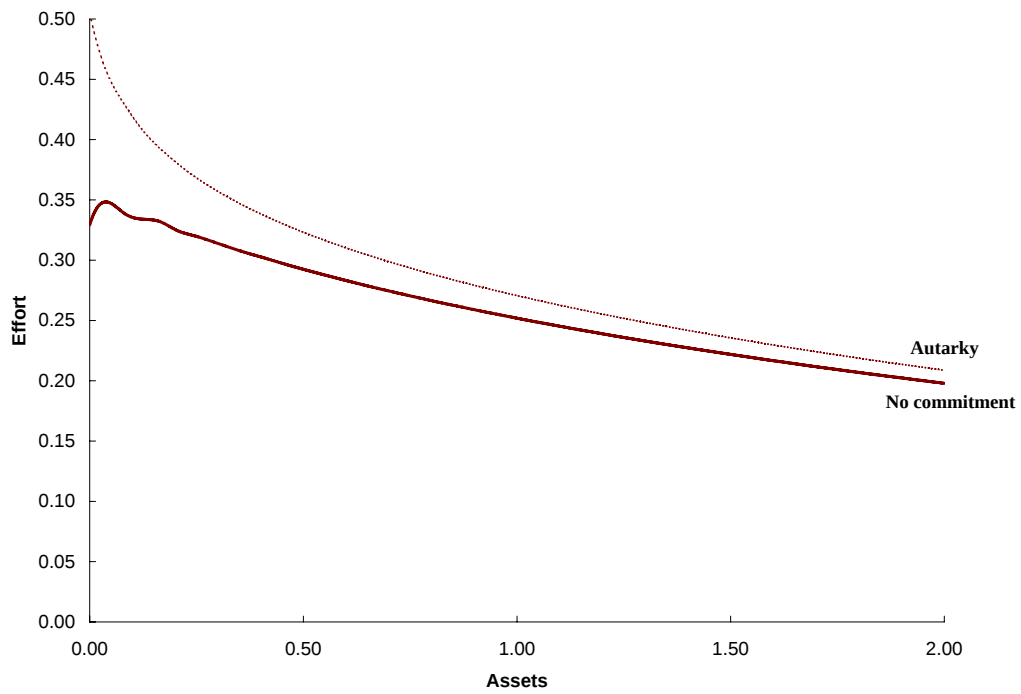


Figure 5: Profit-maximizing principal — No commitment: Expected Discounted profits

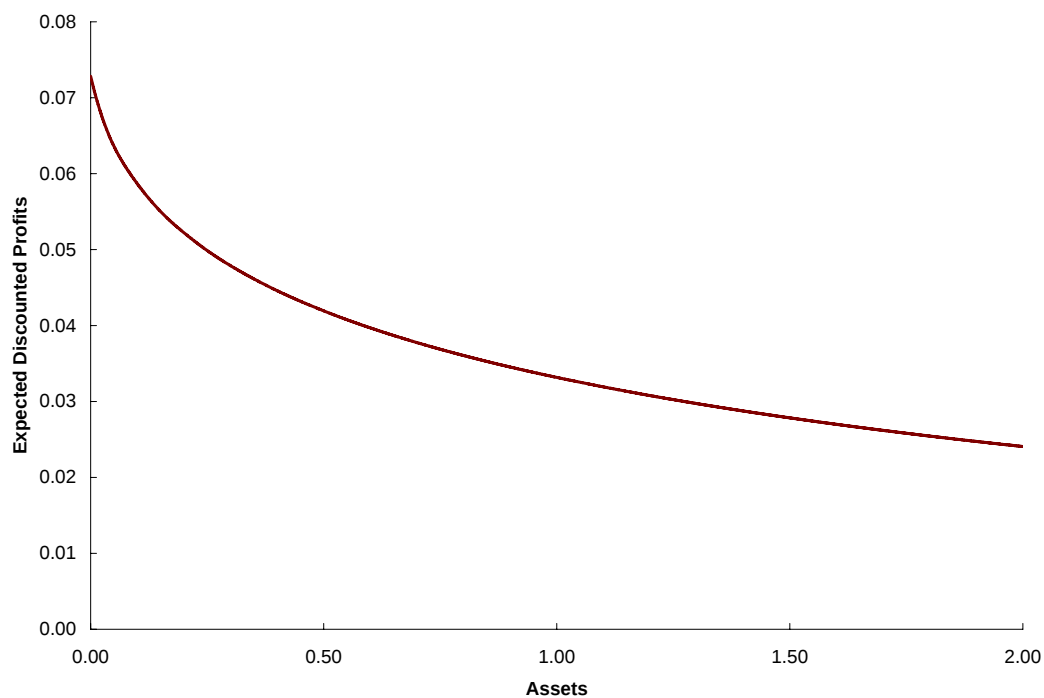


Figure 6: Profit-maximizing principal — No commitment: Profits by case

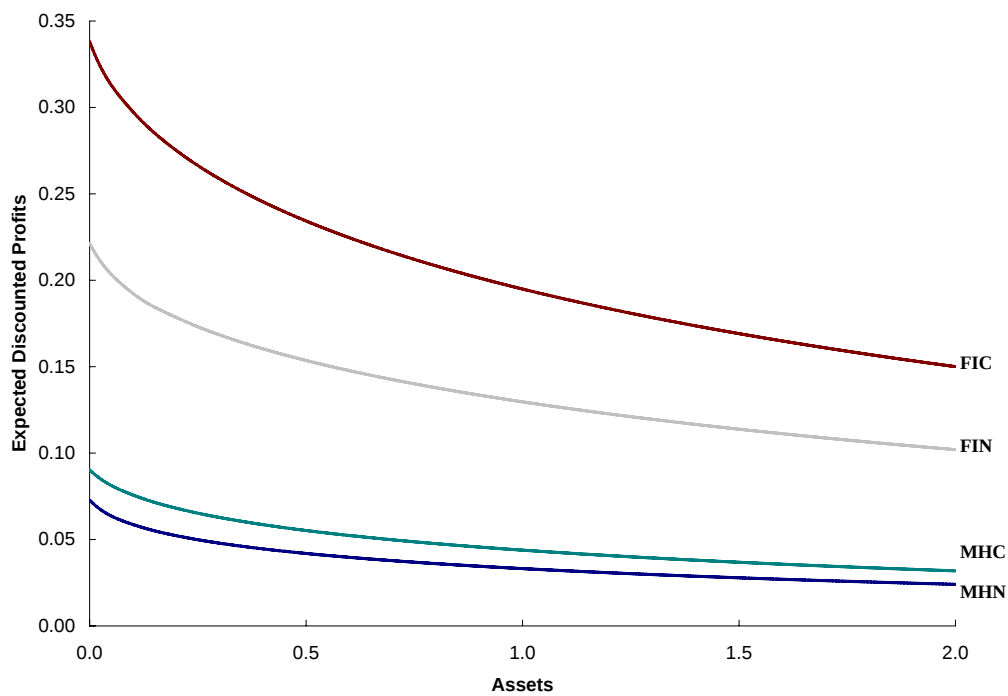


Figure 7: Profit-maximizing principal — No commitment vs. Full commitment: Profits

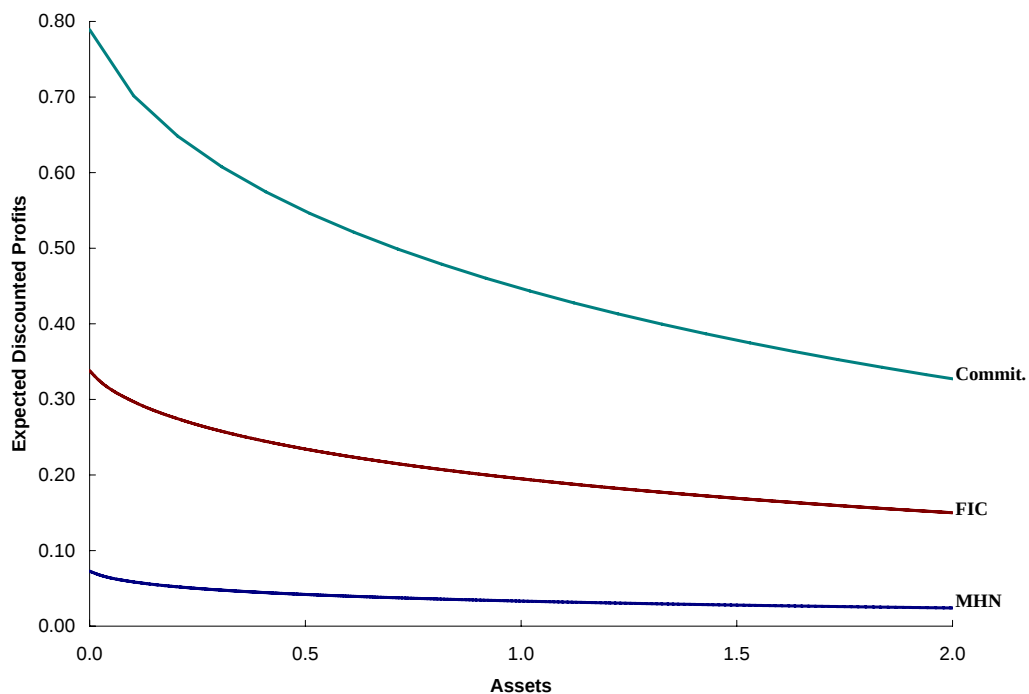


Figure 8: Profit-maximizing principal — No commitment: Lorenz Curves

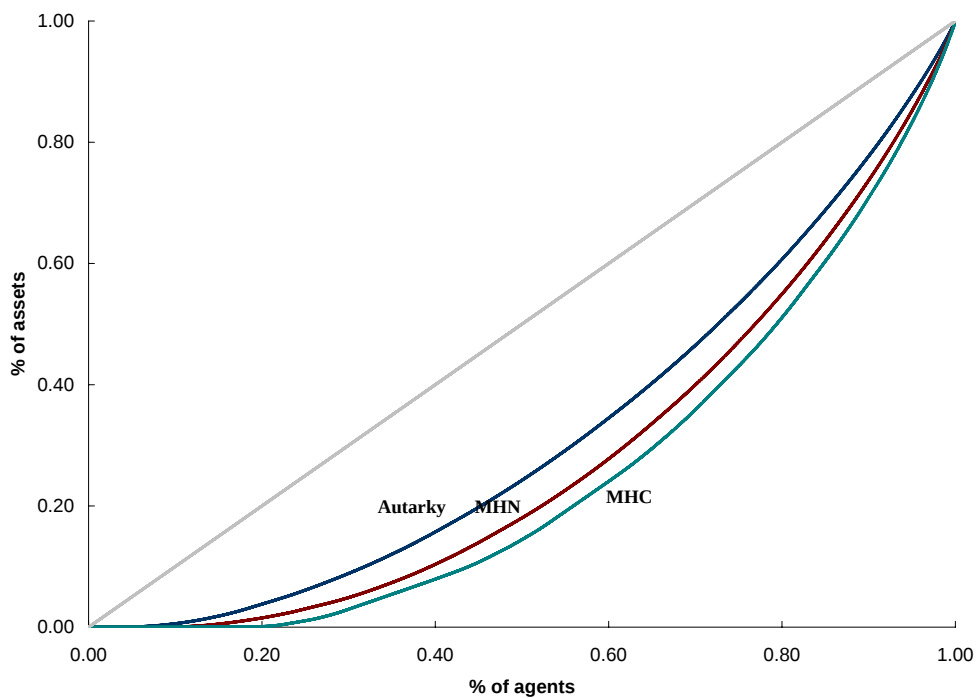


Figure 9: Benevolent planner — No commitment: Consumption equivalence compensation

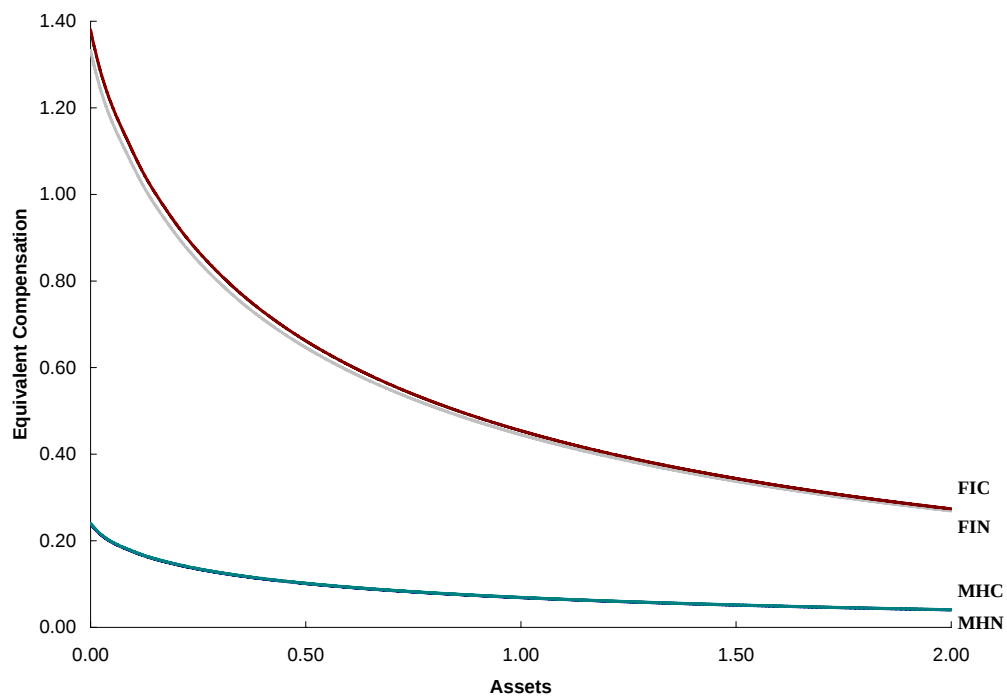


Figure 10: Benevolent planner — No vs. Full commitment: Consumption equivalence compensation

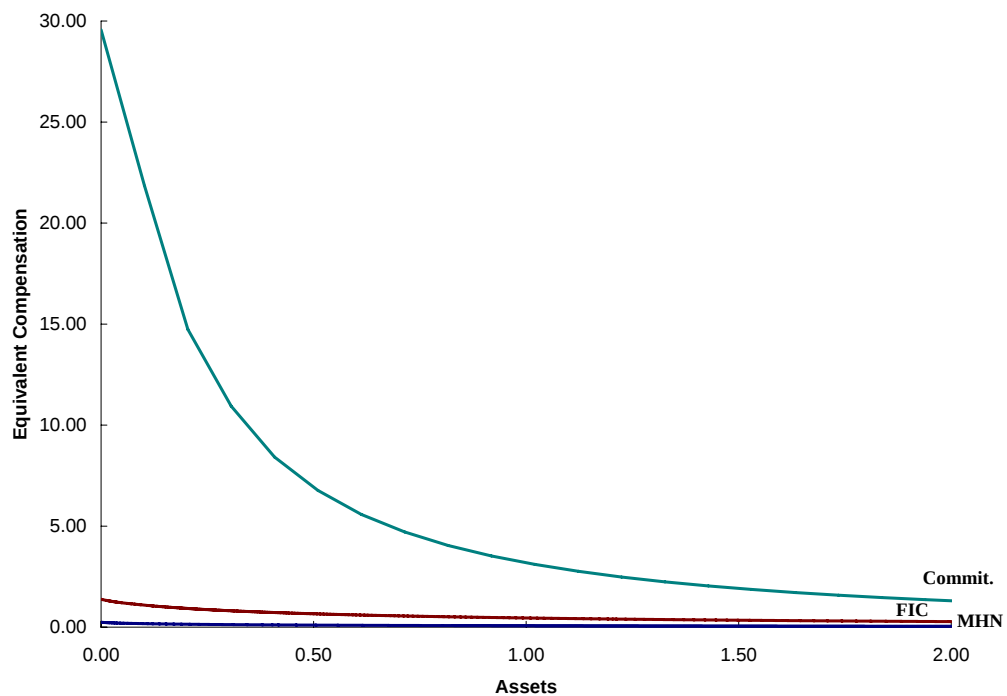


Figure 11: Benevolent planner - No vs. 1-sided commit.: Consumption equivalence compensation

