



DEPARTMENT OF ECONOMICS
DISCUSSION PAPER SERIES

**REINSURING THE POOR: GROUP MICROINSURANCE
DESIGN AND COSTLY STATE VERIFICATION**

Daniel J. Clarke

Number 573
October 2011

Manor Road Building, Oxford OX1 3UQ

REINSURING THE POOR: GROUP MICROINSURANCE DESIGN AND COSTLY STATE VERIFICATION

DANIEL J. CLARKE*

October 11, 2011



Abstract

This paper analyses collusion-proof multilateral insurance contracts between a risk neutral insurer and multiple risk averse agents in an environment of asymmetric costly state verification. Optimal contracts involve the group of agents pooling uncertainty and the insurer acting as reinsurer to the group, auditing and paying a claim only when the group or a sub-group has incurred a large enough aggregate loss. We interpret our models as providing support for insurance contracts between insurance providers, such as microinsurers or governments, and groups of individuals who have access to cheap information about each other, such as extended families or members of close-knit communities. Such formal contracts complement, and could even crowd in, cheap nonmarket insurance arrangements.

Keywords: Microinsurance; Group insurance; Costly state verification; Mechanism design.

JEL codes: D14, D82, G22, O16.

*DPhil student: Department of Economics, Centre for the Study of African Economies ('CSAE') and Balliol College, University of Oxford (clarke@stats.ox.ac.uk; <http://www.stats.ox.ac.uk/~clarke>). This paper forms part of my DPhil thesis, which was supervised by Sujoy Mukerji and Stefan Dercon; the work would not have been possible without their very generous assistance. A number of others have provided very useful comments on aspects of the paper; without implicating them in the shortcomings of the work, I thank my D.Phil. examiners, Christian Gollier and John Quah, as well as Keith Crocker, Ian Jewitt, Markus Loewe, Rocco Macchiavello, Olivier Mahul, Meg Meyer, Joe Perkins, Francis Teal, John Thanassoulis and Liam Wren-Lewis. I have presented the paper at the Gorman Student Research Workshop at the University of Oxford, the CSAE Research Workshop, an ESF SCSS Exploratory Workshop 2009, the CSAE Conference 2010, ERD Regional Conference 2010, the CEAR Insurance for the Poor Workshop 2010, and the FERDI Workshop on Index-Based Weather Insurance 2011. ESRC funding is gratefully acknowledged.

1 INTRODUCTION

Being poor in a poor country is risky. Not only is income from agriculture or the informal sector unpredictable, but the constant threat of health or mortality shocks leaves households vulnerable to serious hardship (Dercon 2004, Collins, Morduch, Rutherford and Ruthven 2009). In the absence of risk management tools the poor remain vulnerable to shocks; a typical response to a large income shock is to take children out of school (Jacoby and Skoufias 1997) and reduce nutritional intake (Behrman and Deolalikar 1988), particularly for girls (Behrman and Deolalikar 1990) and women (Dercon and Krishnan 2000).

One widely observed risk coping mechanism is that of risk pooling within families or communities. Nonmarket risk pooling can be informal (Morduch 2002) or through semiformal institutions that mutualise specific perils such as funerals (Dercon, De Weerd, Bold and Pankhurst 2006), health (Jütting 2004) or fire (Cabral, Calvo-Armengol and Jackson 2003). However, any risk pooling within a small community will be subject to a budget constraint, leaving households vulnerable to shocks that affect the whole community, and may also be constrained by the limited ability of households to commit to state contingent transfers.

Contracting with a formal institution, such as an insurer or government, could break budget and limited commitment constraints but is typically subject to large deadweight transaction costs and information asymmetries. In practice, formal insurance in poor countries is underdeveloped or nonexistent, leaving households exposed to shocks that are not diversifiable within their risk pooling network (Karlan and Morduch 2009, Section 7).

One major cause for the departure from first best risk sharing between formal insurers and poor individuals is the cost of ex-post claims processing, known as *loss adjustment*. Loss adjustment includes both the cost of verifying that claims are not fraudulent (*auditing*, to use the terminology of Townsend (1979))

and the cost of subsequently paying valid claims. In practice, loss adjustment costs are substantial for providers of formal indemnity-based insurance such as reinsurers, insurers and governments.¹ However, loss adjustment may be inexpensive within groups of individuals with intertwined economic or social lives, or within groups of firms in similar or complementary lines of work. Formal loss adjustment can generate a substantial variable deadweight cost and can therefore be a key determinant of the shape, not just the existence, of optimal insurance arrangements.

Much of the existing insurance design literature characterises efficient bilateral contracting between an insurer and a single policyholder. The baseline bilateral insurance contract in this literature is one in which the insurer receives a fixed premium from the policyholder and, in return, offers full marginal insurance below a deductible. Originally motivated by Arrow (1963) as the optimal contract when the insurer set premiums as a defined function of the actuarial value of a contract, it has since appeared as the optimal contract in models of ex ante moral hazard (Hölmstrom 1979) and costly state verification both without (Townsend 1979) and with (Picard 2000) sabotage. There is also a rich literature motivating and explaining more general forms of bilateral insurance contracts.²

Some authors have investigated multilateral insurance contracting. Arnott and Stiglitz (1991) consider the case of optimal contracting between an insurer and multiple agents under ex ante moral hazard where agents can side contract. However they only consider the case where the insurer sells a bilateral contract to each agent and do not consider the form of more general multilateral contracts. Ghatak and Guinnane (1999) and Rai and Sjöström (2004) both consider models of multilateral credit contract design with the former allowing costly state verification and the latter instead allowing agents to send cross reports. However, both assume that agents are risk neutral and therefore ignore any demand for formal insurance.

¹ See for example Derrig (2002), from a special issue of the Journal of Risk and Insurance, focusing on insurance fraud.

² See Dionne, ed (2000) for an introduction to the literature.

This paper considers optimal multilateral contracting between a risk neutral insurer and two risk averse agents whose losses are affiliated and where the cost of loss adjustment, that is auditing and making transfers, is higher for the insurer than for agents. Auditing is modelled in Townsend's (1979) deterministic costly state verification framework; agents are required to send messages to the insurer, who only learns the true loss incurred by an agent by conducting an audit. In addition to not auditing or fully auditing the state of the world, the insurer may partially audit the state by auditing only one loss, rendering the problem a nontrivial multidimensional extension to Townsend (1979). The deadweight cost of loss adjustment by the insurer is assumed to increase in both the claim payment and the number of audits, and extends Arrow's (1963) single argument actuarial value cost function to two dimensions. By contrast, agents learn each other's losses and may make side transfers to each other at zero cost.

Agents are allowed to collude against the insurer and so, as is common in bilateral models of deterministic costly state verification, the aggregate net transfer from insurer to agents depends only on losses audited by the insurer; were the contract to depend on unaudited information, agents would always collude to send the jointly most favourable message, sharing the gains between themselves.

In the benchmark model we also allow each agents to increase their ex post loss by *sabotage*, without detection by the insurer. Allowing agents to conduct sabotage constrains the optimal contracts to feature no marginal overinsurance, since there will always be a no-sabotage contract that dominates any contract with sabotage.

The economic problem is to design an arrangement that utilises the loss adjustment technology in an efficient fashion. Nonmarket insurance between agents is inexpensive but subject to a budget constraint. Formal insurance is not subject to a budget constraint but is expensive. As one might expect, efficient arrangements involve the agents offering mutual nonmarket insurance and the insurer offering protection for losses that are large for the group, relative to the corresponding increase in loss adjustment cost. The formal insurer therefore acts

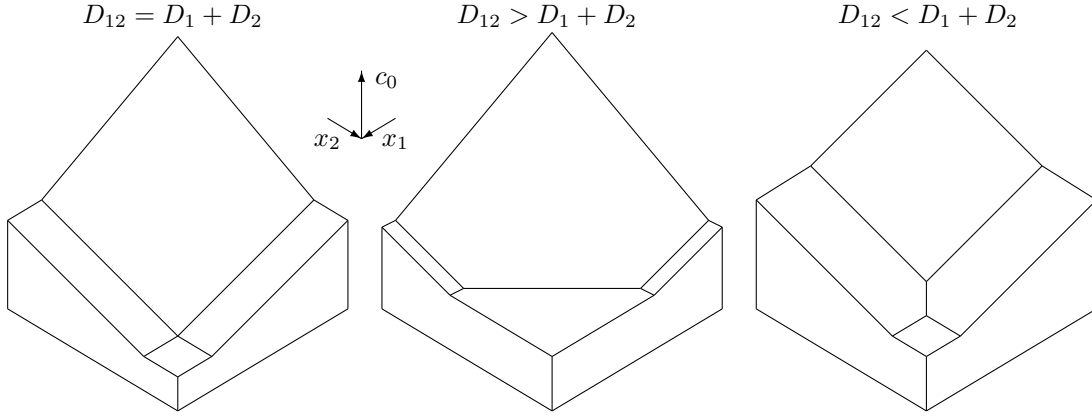
as reinsurer to the group who, in turn, can pool uncertainty between themselves at low informational and transactional cost.

The optimal solution in our benchmark model takes a particularly simple form, which we term a *Generalised Stop Loss* contract, named after the *Pure Stop Loss* contracts common in markets for reinsurance. For a group of two agents, an aggregate premium of p_0 is paid to the insurer, who in turn makes aggregate claim payment of $\max(0, x_1 - D_1, x_2 - D_2, x_1 + x_2 - D_{12})$ where x_1, x_2 are each agent's loss and D_1, D_2, D_{12} are single and double loss deductibles. This is the natural multidimensional extension to the bilateral insurance contract which offers full marginal insurance below a deductible.

Figure 1 shows an isometric projection of the total consumption of two agents who have jointly purchased a symmetric Generalised Stop Loss contract. The aggregate consumption schedule in the left pane could be achieved by both agents purchasing bilateral contracts from an insurer, each offering full marginal insurance against loss i above a deductible of D_i . However the consumption schedules in the middle and right pane could only be achieved by a contract where the claim payment to at least one of the agents is conditional on both losses. The key requirements for multilateral contracting to dominate bilateral contracting is that the group has access to cheap within-group loss adjustment technology, enabling cheap within-group monitoring and transfers. As is shown in section 6, formal insurers can use their contracting power to crowd in risk pooling within groups, and so within-group limited commitment or enforcement constraints do not invalidate the group-based approach to insurance.

When the insurer is able to detect sabotage for one loss, thereby removing the restriction of no marginal overinsurance for that loss, we find that optimal contracts are similar to Generalised Stop Loss contracts but with marginal overinsurance on the region for which only that loss is audited. The marginal overinsurance is optimal since a high loss incurred by one agent is associated with a high loss incurred by the second agent. This optimal contract is similar

Figure 1. Aggregate net consumption of agents under a symmetric Generalised Stop Loss contract with single loss deductibles of $D_1 = D_2$ and double loss deductible of D_{12}



to that of *model plot* area yield index agricultural insurance, in which the claim payout to all insured farmers within a locality depends on the average audited yield from a sample of model plots, chosen to be manipulation-free and statistically representative of the locality (Miranda 1991, Mahul 1999). In our framework such statistical sample-based indices are only ever optimal if the insurer can verify whether sabotage has occurred in the model plots. Moreover, the technology that allows the insurer to verify whether no sabotage has occurred is valuable, in that it allows a multilateral insurance contract to be put in place that strictly dominates all Generalised Stop Loss contracts.

The final model investigate the optimal contract when the insurer can condition claims on a costlessly observable index as well as auditable losses. The optimal contract is similar to the Index Plus Gap insurance policies of Doherty and Richter (2002), in which the indemnity-based element offers marginal insurance against losses net of indexed claim payments. In our multidimensional framework we find that the insurer pays an indexed payment and that the indemnity-based gap insurance is of the Generalised Stop Loss form, again based on losses net of the index claim payment. Pure index insurance is not optimal in general due to the possibility of the index realisation being good, but the aggregate group losses being high. In such joint states of the world, it may be optimal for loss adjustment

to be conducted and a claim payment made Clarke (2011).

These models offer a template for insurers wishing to offer efficient insurance arrangements to poor individuals. First, they should consider contracting with groups of individuals, economically and socially close enough to have access to a cheap loss adjustment technology and small enough to be able to sustain risk pooling. Second they should use their contracting power to support nonmarket insurance between group members wherever possible. Third, they should think of their role as that of a reinsurer, offering protection when the group as a whole is unlucky. They should not offer cover for idiosyncratic shocks, as this would crowd out cheap within-group pooling. Fourth, both cheaply available indices that are affiliated with losses, and audit technologies that allow an insurer to verify when sabotage has taken place are valuable for the insurer, since they allow Pareto dominant insurance contracts to be signed between the insurer and agents.

Stop Loss Reinsurance in the Small and in the Large

These principles are already observable in insurance contracting around the world.

One example of pooling plus Stop Loss reinsurance is that of the Self-Insurance Funds (Fondos de Aseguramiento Agrícola) that have been operating in Mexico since 1988. Under the standardised legal framework farmers may join or create a Self-Insurance Fund, which allows farmers to pool agricultural uncertainty, but any such Fund must purchase Stop Loss reinsurance from a licensed insurance company. In 2004 Agroasemex, a reinsurer owned by the Mexican Government, sold pure Stop Loss products to 242 groups with total membership of more than 70,000 farmers (Ibarra 2004). The reinsurance product indemnifies each Self-Insurance Fund against total Fund losses above a Fund-level deductible. Unusually, it is a voluntary indemnity-based crop insurance scheme that has been self-funding in the medium term: as reported by Ibarra (2004), the total insurance

payout by Agroasemex since inception was 57.7% of the total premium income by Agroasemex, above the break even net loss ratio of 75%.

The financial structure suggested in this paper for insurers and governments contracting with poor individuals is surprisingly similar to that observed between reinsurers and insurers or reinsurers and captive insurance companies in developed financial markets (Swiss Re 2002). One particularly famous arrangement that has existed since at least the 18th Century is the global structure of protection and indemnity (P&I) insurance for shipowners (Bennett 2000). The thirteen major P&I Clubs and the International Group of P&I clubs (IGP&I) coordinates not-for-profit mutual liability insurance for approximately 90% of the world's ocean-going tonnage. The IGP&I then purchases one reinsurance policy from international reinsurers to protect shipowners worldwide against the IGP&I incurring a very large loss, which each shipowner is jointly and severally liable for. The theory of this paper provides a positive interpretation of such a structure if P&I Clubs, whose members have close economic ties, can conduct loss adjustment at a lower cost than the external reinsurer.

The rest of the paper is organised as follows. Sections 2 and 3 set up and characterise optimal contracts in our benchmark model. In Section 4 we allow the insurer to discriminate between sabotage and raw losses for one agent, in Section 5 we allow the insurance arrangement to be conditioned on a costlessly observable index, and in Section 6 we consider the case where agents cannot commit to a complete state contingent contract. Section 7 concludes. All proofs, unless otherwise noted in the text, are in the Appendix.

2 THE BENCHMARK MODEL

2.1 Preferences, Losses, and Information

Two agents contract with a single insurer. Preferences for agent i are represented by a twice differentiable utility function u_i , defined over own consumption c_i . Agents are strictly non-satiated and risk-averse, that is $u'_i > 0$ and $u''_i < 0$ everywhere for all agents i . The insurer is risk-neutral, maximising expected profits.

Each agent i has initial wealth w_i and is subject to uncertain loss x_i which can take values in the interval $[0, \bar{x}]$. A state of the world x is a pair of losses $x = (x_1, x_2) \in X$ where $X = [0, \bar{x}] \times [0, \bar{x}]$. The exogenous cumulative density function of x is common knowledge and denoted by $F(x)$. F is atomless with full support on X , and losses are *affiliated* in the sense of Milgrom and Weber (1982); that is we can denote the probability density function by $f : X \rightarrow (0, \infty)$ where

$$f(x)f(x') \leq f(x \vee x')f(x \wedge x') \text{ for all } x, x' \in X \quad (1)$$

with $x \wedge x' = (\min\{x_i, x'_i\})_{i=1}^2$ and $x \vee x' = (\max\{x_i, x'_i\})_{i=1}^2$. Loss affiliation captures the notion that losses x_i and x_j are everywhere positively correlated, and will act to ensure that the optimal claim payment when only one loss is audited is nondecreasing in the loss.

Each agent i observes the realised joint state of the world x costlessly and may then increase own loss x_i by sabotage of $s_i \in [0, \bar{x} - x_i]$. Sabotage choices are observed by the other agent and then each agent sends a message m_i to the insurer. The insurer can observe zero, one or both losses by conducting costly audits. Following Picard (2000), the audit technology does not allow the insurer to distinguish between raw loss x_i and loss due to sabotage s_i , but is otherwise

perfect; if the insurer audits agent i ($a_i = 1$), it discovers the total loss $x_i + s_i$; if the insurer doesn't audit agent i ($a_i = 0$), it discovers nothing.³ We will find that allowing agents to conduct sabotage ensures that optimal contracts will feature no marginal overinsurance. (We relax the assumption that the insurer cannot discriminate between raw loss x and sabotage s in Section 4.)

Each agent i makes net transfer θ_i to the insurer and net side transfer τ_i to the other agent. Side transfers are not observable by the insurer. The resulting ex-post consumption of agent i is denoted by:

$$c_i = w_i - x_i - s_i - \theta_i - \tau_i \quad (2)$$

Loss adjustment, that is auditing and making transfers θ_i , incurs a deadweight cost for the insurer. We defer full specification of the deadweight loss adjustment cost function $\kappa(\cdot)$ until we have introduced and rewritten the incentive compatibility constraints in section 2.3.

2.2 Mechanisms and Side Contracts

The ex-post asymmetry in the cost of loss adjustment leads to an implementation problem similar to that considered by Townsend (1979). However, while Townsend (1979) considered the form of optimal contracts between an insurer and a single agent, we will consider the form of an optimal multilateral mechanism between an insurer and multiple agents.

A mechanism $G = (\{M_i, a_i, \theta_i\}_{i=1,2})$ will be offered to agents and the insurer by an independent mechanism designer, where each element has the following interpretation: M_i is the message space of agent i ; $a_i : M_1 \times M_2 \rightarrow \{0, 1\}$ is the auditing function (audit (1) or no audit (0)); and $\theta_i : M_1 \times M_2 \times X \rightarrow \mathbb{R}$ is the net

³ Our sabotage assumption may be considered as an extreme form of costly state falsification, where an agent can falsify a loss x_i as $x'_i \geq x_i$ at cost $x'_i - x_i$. Optimal insurance contracts under costly state falsification, where the marginal cost of falsification is less than unity can feature overpayment of small and underpayment of large claims (Bond and Crocker 1997, Crocker and Morgan 1998).

transfer from agent i to the insurer, and is restricted to vary only with messages and audited losses, not unaudited losses.⁴ As is common in insurance models with costly state verification, the insurer is restricted to choose an audit rule that is a deterministic function of joint message m .⁵ If either agent or the insurer reject G the insurer receives zero profit and each agent i receives reservation utility $\bar{u}_i$.

The form of the optimal mechanism will depend on the ability of agents to write binding side contracts with each other. If agents cannot contract with each other at all over messages m , side-transfers τ or the sabotage of losses s then the insurer can implement full insurance without any auditing in equilibrium:

Proposition 1. *If agents cannot side-contract then any consumption schedule $(c_1(x), c_2(x))$ such that each c_i is nonincreasing in x_i can be implemented by dominant strategies, with no auditing or sabotage in the equilibrium. If each c_i is strictly decreasing in x_i for all x_i then the implementation is unique.*

The insurer is able to fully extract information by setting up a Prisoner's Dilemma game to be played by the agents, where the insurer audits both agents whenever reports of the total loss disagree, rewards any truthful agent, and punishes untruthful agents. So long as each $c_i(x)$ is nonincreasing in own net loss then each agent will have no incentive to conduct sabotage, and if $c_i(x)$ is strictly decreasing in own loss then sabotage is never optimal.

If there is no deadweight loss adjustment cost when no audits are conducted, Proposition 1 implies that the first best, where full insurance is offered to both risk averse agents with no deadweight loss adjustment cost, can be implemented. Moreover, unique implementation is possible for a schedule that approximates full insurance.

⁴ Formally, the restriction on θ_i is $\theta_i(m_1, m_2, x) = \theta_i(m_1, m_2, x')$ for all $(m_1, m_2) \in M$, $x, x' \in X$ such that $a_k(m_1, m_2) \times (x_k - x'_k) = 0$ for $k = 1, 2$.

⁵ An exception is Fagart and Picard (1999), who characterise the optimal bilateral insurance contract under costly state verification when the audit rule is stochastic. However, in their model the global incentive compatibility constraints do not reduce to a local first order constraint except in the special case where the risk averse agent has constant absolute risk aversion. In another work Krasa and Villamil (2000) are able to show that optimal stochastic audit rules reduce to deterministic audit rules in a model of costly state verification with a form of renegotiation-proofness.

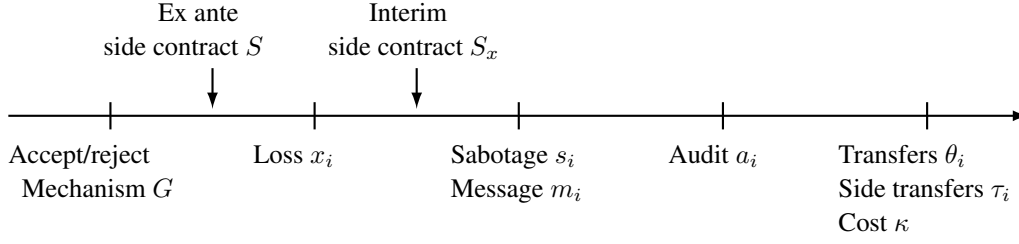
However, the assumption that agents cannot side contract at all seems unrealistic, particularly in situations where agents collectively have much to gain from conspiring against the insurer. We will henceforth assume that some degree of coordination between agents is possible. Such coordination could be modelled by an extensive form bargaining game between the agents over messages to be sent and corresponding side-transfers between agents. Following Laffont and Martimort (2000) and Rai and Sjoström (2004) we abstract from this approach, and allow agents to sign binding side-contracts. We assume that the side-contracts are enforceable, although we do not specify a court of justice able to enforce such contracts.⁶ Any enforcement constraints will be modelled in reduced form as restrictions on the side-contracts that may be signed. This is a modelling short cut, which allows us to capture in a static context the reputations of the agents which would guarantee self-enforceability in a repeated relationship. In the context of microinsurance, agents are likely to be able to threaten social sanctions, exclusion from nonmarket insurance or physical violence.

A side contract is a specification of the sabotage s_i to be performed by each agent, the message m_i to be sent by each agent to the insurer, and the net side-transfer τ_i from agent i to j , satisfying $\tau_1 + \tau_2 \geq 0$. By allowing agents to side contract, the mechanism designer cannot hope to elicit information by rewarding an individual agent, as that agent can commit to transfer any such reward to the other agent. The mechanism of Proposition 1 would unravel, with agents jointly choosing sabotage amounts and messages with the lowest aggregate transfer to the insurer, $\theta_0 = \theta_1 + \theta_2$, and the insurer never detecting the fraud.

Following Rai and Sjoström (2004) we will consider two extreme classes of side contracts which differ in the ability of agents to pool uncertainty by committing to state contingent side transfers. A side contract where agents can fully commit to a state contingent transfer rule before the state of the world x is realised, is an *ex ante side contract* $S = (\{s_i, m_i, \tau_i\}_{i=1,2})$ where $s_i : X \rightarrow X$, $m_i : X \rightarrow$

⁶ However, note that in many semiformal risk pooling arrangements there is an explicit process for arbitration and enforcement (Dercon et al. 2006, Jütting 2004, Cabrales et al. 2003).

Figure 2. Timeline



M_i and $\tau_i : X \rightarrow \mathbb{R}$. Under ex ante side contracting agents are able to sign Pareto optimal side contracts, performing sabotage and sending messages that maximise joint consumption $c_0 = c_1 + c_2$ and committing to a Pareto optimal consumption sharing rule. We consider the case of *interim side contracting* in Section 6, where agents are only able to sign binding contracts after the state of the world is realised. Under interim side contracting agents are able to collude against the insurer but are not able to commit to a Pareto optimal consumption sharing rule.

The timing is shown in figure 2.

2.3 Feasibility

Some restrictions on mechanisms are now in order. The mechanism designer will restrict attention to mechanisms which satisfy certain incentive compatibility and individual rationality (participation) constraints. Mechanism G and ex ante side contract S are together individually rational if expected utility for each agent exceeds the respective reservation utility, and the insurer receives nonnegative expected profit. Moreover, for a given mechanism G , an ex ante side contract S will be called incentive compatible if the agents could not sign a side contract that would be better for both agents:

(IR1) (For individual rationality under ex ante side contracting)

$$\mathbb{E}[\theta_0] - \kappa \geq 0 \text{ and } \mathbb{E}[u_i(c_i)] \geq \bar{u}_i \text{ for } i = 1, 2.$$

(IC1) (For incentive compatibility of an ex ante side contract)

There is no other individually rational side contract which gives strictly higher expected utilities to both agents.

Under ex ante side contracting the revelation principle holds and allows us to narrow focus to direct mechanisms, where both agents report the joint total loss $x + s \in X$. We represent a direct mechanism as $G = (a, \theta)$, suppressing the message space $M = M_1 \times M_2$ which is taken to be $X \times X$. We will denote the outcomes when messages are truthful as indirect functions of the total loss $x + s$. That is, $a(x + s) := a(x + s, x + s)$, $\theta(x + s) := \theta(x + s, x + s, x + s)$ and $c(x + s) := w - x - s - \theta(x + s) - \tau(x + s)$ for all $x \in X$. For a given side-contract, the restrictions on mechanisms that guarantee that truth-telling messages are optimal for the agents will be called *incentive compatibility*:

(IC2) (For truthful report of the net loss)

There is an incentive compatible side contract with truthful messages, $m_1(x + s) = m_2(x + s) = x + s$ for all $x + s \in X$.

A mechanism G will be called *feasible* if there exists some ex ante side contract S which, together with G , satisfies (IR1), (IC1) and (IC2). For a feasible mechanism we will assume that all parties accept the mechanism and agents sign a side contract which, together with the mechanism, is individually rational and incentive compatible with truthful message rule $m(x) = (x, x)$ for all $x \in X$.

Optimality of a mechanism will be taken to mean constrained Pareto dominance in the following sense. A mechanism will be defined to *weakly dominate* another mechanism if for every feasible side contract for the latter mechanism, there is a feasible side contract for the former mechanism such that the expected utilities for all agents and the insurer are at least as large as in the other mechanism and side contract. The former *dominates* the latter if the weak dominance relation holds with strict inequality for at least one agent. A mechanism is defined to be optimal if it is not dominated by any other mechanisms. As usual, the use of

Pareto dominance in the definition of optimality allows us to abstract from the question of how the gains from trade should be split. As will become clear, the optimal form of the contract will not depend on the allocation of gains from trade.

We may now begin to characterise optimal feasible insurance mechanisms. First we will show that any feasible mechanism is weakly dominated by a mechanism with no marginal overinsurance, that is where

$$x_0 + \theta_0(x) \text{ is weakly increasing in each } x_i, \ i = 1, 2, \quad (3)$$

for $x_0 = x_1 + x_2$.

Lemma 1. *Any direct feasible mechanism under ex ante side contracting (a, θ) is weakly dominated by a direct feasible mechanism (a', θ') that satisfies condition (3)*

Under a feasible mechanism that satisfies condition (3) and for any incentive compatible side contract $S = (s, m, \tau)$ then $s(x) = 0$ almost everywhere.⁷ Therefore without loss of generality we restrict attention to feasible mechanisms that satisfy condition (3) and incentive compatible side contracts in which $s(x) = (0, 0)$ for all $x \in X$.

Second we note that θ_0 can only vary with audited information, that is losses that the insurer has learned through auditing. If this were not the case, and θ_0 varied with messages sent by the agents even as the audited information remained the same, then any side contract that satisfied (IC1) would specify that in each state of the world agents would send the joint message that resulted in the lowest $\theta_0 = \theta_1 + \theta_2$, violating (IC2). We may therefore state the following result which extend's Townsend's (1979) rewrite of the bilateral incentive compatibility constraints to the multilateral case.

Lemma 2. *If $a(x)$ and $\theta_0(x)$ are feasible then there is a constant p_0 and functions*

⁷ That is to say the set $\{x \in X | s_1(x) + s_2(x) > 0\}$ has zero measure.

$y_1(x_1)$, $y_2(x_2)$ and $z(x)$ such that

$$\theta_0(x) = p_0 - \max(y_1(x_1), y_2(x_2)) - z(x) \quad (4a)$$

$$a_i(x) = 1 \text{ if } \begin{cases} y_i(x_i) > y_j(x_j) \text{ or} \\ y_i(x_i) = y_j(x_j) > 0 \text{ and } a_j(x) = 0 \text{ or} \\ z(x) > 0 \end{cases} \quad (4b)$$

$$\text{where } p_0, y_1(x_1), y_2(x_2), z(x) \geq 0 \text{ for all } x \in X. \quad (4c)$$

The notation may be interpreted as follows. p_0 may be interpreted as the total premium paid by group members to the insurer, with $p_0 - \theta_0(x)$ as the total claim payment to group members in state x . $y_1(x_1)$ and $y_2(x_2)$ are minimum claim payments payable to the group based on either agent 1's or agent 2's losses respectively, such that if the state is x the total claim payment to the group is at least $\max(y_1(x_1), y_2(x_2))$. $z(x)$ is the additional claim payment payable in addition to $\max(y_1(x_1), y_2(x_2))$ if a second loss is audited.

A sketch of the proof is as follows. First, as is typical in models with deterministic auditing rules, θ_0 must be constant over the region where no audits occur. If not, the group would have an incentive to side contract so that equilibrium reports were of the zero audit state of the world with the lowest θ_0 . Truthtelling would not be incentive compatible, violating (IC2). We may therefore write $\theta_0(x)$ for those x such that $a_1(x) = a_2(x) = 0$ as some constant p_0 .

Second, unaudited states cannot result in a higher claim payment from insurer to agents than audited states, or the group would have an incentive to report a zero audit state, violating (IC2). Therefore $\theta_0(x) - p_0$ must be nonnegative for all x . Since expected insurer profits must be nonnegative, $p_0 \geq 0$.

Third, consider some x_1, x'_1, x_2, x'_2 such that in state (x_1, x'_2) only loss 1 is audited and in state (x'_1, x_2) only loss 2 is audited. For the group not to have

the incentive to misreport either (x_1, x'_2) or (x'_1, x_2) when the true state was (x_1, x_2) it must be that $\theta_0(x_1, x_2) \leq \theta_0(x_1, x'_2)$ and $\theta_0(x_1, x_2) \leq \theta_0(x'_1, x_2)$. When only one agent is audited in state (x_1, x_2) it must be that $\theta_0(x_1, x_2) = \min(\theta_0(x_1, x'_2), \theta_0(x'_1, x_2))$. For each x_i , if there is some x_j such that only agent i is audited, we define $y_i(x_i) := p_0 - \theta_0(x_1, x_2)$ and if there is no such x_j we define $y_i(x_i) := 0$. When only one agent is audited we may therefore write $\theta_0(x) = p_0 - \max(y_1(x_1), y_2(x_2))$.

Fourth, in any state x where both agents are audited the above inequality need not bind, and so we have $\theta_0(x) \leq p_0 - \max(y_1(x_1), y_2(x_2))$. Therefore there exists a $z(x) \geq 0$ such that (4a) holds.

2.4 Loss Adjustment Cost

We are now able to state our assumption about the cost of loss adjustment. Define $y(x) := \max(y_1(x_1), y_2(x_2))$. Then:

Assumption 1. *The deadweight loss adjustment cost to the insurer of a feasible mechanism $\{a, \theta\}$ is $\kappa(\mathbb{E}y, \mathbb{E}z)$ where $\kappa(0, 0) \geq 0$ and $D_2\kappa(Y, Z) \geq D_1\kappa(Y, Z) > 0$ for all $Y, Z \geq 0$.⁸*

Assumption 1 is in effect an extension of Arrow's (1963) assumption that the insurer is willing to offer an insurance policy at a premium which depends only on the policy's actuarial value, that is the expected claim payment. In our environment the expected deadweight cost depends on the auditing structure as well as the actuarial value. The assumption $D_2\kappa \geq D_1\kappa > 0$ captures the notion that loss adjustment is costly and that it is at least as costly to increase the claim payment when two audits are necessary as when only one audit is necessary.

Suppose that the mechanism designer wanted to increase the claim payment in some state x where $y_i(x_i) > y_j(x_j)$. $y_i(x_i)$ could be increased, but this would necessitate increasing the minimum claim payment for all states x' in which agent

⁸ The notation $D_i\kappa$ denotes the partial derivative of κ with respect to the i th argument.

i 's loss was x_i . The designer could instead increase $z(x)$, without the need to increase claim payments in other states, but this increase would be subject to a larger deadweight cost than that from increasing $y(x)$ in state x .

Arrow's (1963) assumption is a special case of assumption 1, where $\kappa(\mathbb{E}y, \mathbb{E}z) = \kappa(\mathbb{E}y + \mathbb{E}z)$. A simple linear example of a loss adjustment function that satisfies assumption 1 is where the realised cost of loss adjustment in some state x is $\kappa_0 + \kappa_1 \times y(x) + \kappa_2 \times z(x)$ for constants $\kappa_2 \geq \kappa_1 > 0$ and $\kappa_0 \geq 0$.

3 THE BENCHMARK MODEL AND STOP LOSS CONTRACTS

Under ex ante side contracting we first reduce the multilateral contracting problem between two agents and an insurer to a bilateral problem between an insurer and a group representative agent. We parameterise optimal contracts by λ , the expected profit of the insurer, and show that any optimal mechanism maximises the expected utility of a group representative agent, subject to our rewritten incentive compatibility constraints. Denoting the expected profit of the insurer by $\pi(p_0, \mathbb{E}y, \mathbb{E}z) := p_0 - \mathbb{E}y - \mathbb{E}z - \kappa(\mathbb{E}y, \mathbb{E}z)$ and the group consumption by $c_0 := w_0 - p_0 - x_0 + \max(y_1, y_2) + z$ we have the following:

Lemma 3. *If a mechanism (a, θ) is optimal under ex ante side contracting then there exists constant λ and strictly increasing concave function u_0 such that $\theta_0 = p_0 - \max(y_1, y_2) - z$ is a solution to:*

$$\max_{p_0, y_1, y_2, z} \mathbb{E}u_0(c_0) \text{ subject to } \pi = \lambda, (3) \text{ and } (4c) \quad (5)$$

Following Wilson (1968) we may interpret u_0 as a surrogate group utility function, defined over group consumption $c_0 = c_1 + c_2$, with corresponding

belief $f(x)$. The objective function is therefore the expected utility of the group representative agent, and is an increasing linear function of the both realised utility of agent 1 and that of agent 2 under a Pareto optimal sharing rule.

Program (5) implicitly specifies a_i for each loss i of the state x , but only aggregate transfers θ_0 between insurer and group are specified: the program does not specify the split of insurance claim payments between agents. The specific allocation of claim payments does not matter because agents can contract with each other to undo any allocation of the consumption good determined by the insurer. Indeed, agents will contract so that net consumption profiles on the truthful equilibrium path satisfy the Borch (1962) rule.⁹ Neither does the program specify the punishments levied on agents out of equilibrium, but without loss of generality we may assume that when any message is found to be incorrect, the aggregate net transfer to insurer from agents is p_0 .

Our reduced problem is mathematically similar to those considered in Townsend (1979), but with the addition that the state can be partially audited; the insurer can choose to audit just one loss, and therefore only partially learn the state of the world, whereas in Townsend's models the state was one-dimensional and so any audit yielded the full state.

A mechanism (a, θ) is said to be a Generalised Stop Loss contract if

$$\theta_0(x) = p_0 - \max(0, x_1 - D_1, x_2 - D_2, x_0 - D_{12}) \text{ for almost all } x \in X \quad (6)$$

for some $D_1, D_2 \in [0, \bar{x}]$ and $D_{12} \in [0, 2\bar{x}]$. A Generalised Stop Loss contract offers full marginal insurance for any one loss above a single loss deductible of D_i , and full marginal insurance for the combined loss above a group loss deductible of D_{12} . Net consumption of the pair of agents for almost all $x \in X$ is:

$$c_0(x) = w_0 - p_0 - \min(x_1 + x_2, D_1 + x_2, x_1 + D_2, D_{12}) \quad (7)$$

⁹ The Borch Rule states that $u'_1(c_1(x))/u'_2(c_2(x))$ is constant for almost all $x \in X$. If not, there would be a side-contract with truthful equilibrium consumption that strictly dominated the original side-contract.

Figure 1 illustrates the net consumption under a Generalised Stop Loss contract for different individual and group deductibles D_1 , D_2 and D_{12} .

We are now in a position to state our main result.

Theorem 1. *Any optimal feasible mechanism is a Generalised Stop Loss contract.*

The solution methods we employ to prove Theorem 1 follow Winton (1995) and Gollier and Schlesinger (1996) in relying only on notions of second order stochastic dominance. The two key assumptions underlying our result are our loss adjustment cost assumption and the assumption that losses are affiliated. Without the affiliation assumption we could not guarantee that indemnity schedules would be increasing in the incurred loss.

An alternative loss adjustment cost assumption would be that the realised loss adjustment cost depended only on the number of audits conducted. However, this extension of Townsend's (1979) constant fixed audit cost assumption is not easy to work with in a model with incentive compatibility restrictions (4a) and (3). For example, Picard (2000) considered optimal contracting under these assumptions in the bilateral case but could only characterise the optimal indemnity schedule after having restricted attention to schedules in which the claims payment was weakly increasing in incurred loss. In our model Theorem 1 holds under the fixed claim cost assumption if the pdf of losses f is weakly decreasing in both its arguments, but the optimal contract seems more difficult to characterise under less restrictive assumptions.

Theorem 1 follows through if agents are unable to sign complete state contingent contracts on (s, m, θ) but may still commit to a sharing rule, $(c_1(c_0), c_2(c_0))$ as the sharing rule would align the agent's incentives for sabotage and message decisions. We will further relax the assumption about the ability of agents to commit to ex ante side contracts in section 6 and will consider optimal contracts

when sabotage is not possible, or is separately observable by the insurer, in section 4.

Theorem 1 also extends to a setting with N agents: the insurer would offer full marginal insurance for losses above a suite of individual, subgroup and group deductibles.

We may derive a corollary to our main Theorem by appealing to the loss adjustment function considered by Arrow (1963). A mechanism (a, θ) is said to be a Pure Stop Loss contract if

$$\theta_0(x) = p_0 - \max(0, x_0 - D_{12}) \text{ for almost all } x \in X \quad (8)$$

for some $D_{12} \in [0, 2\bar{x}]$. A Pure Stop Loss contract offers full marginal insurance for combined losses above the group deductible D_{12} . It is a special case of a Generalised Stop Loss contract with $D_1 = D_2 = \bar{x}$ or $D_{12} < \min(D_1, D_2)$. Net consumption of the pair of agents for almost all $x \in X$ is:

$$c_0(x) = w_0 - p_0 - \min(x_1 + x_2, D_{12}) \quad (9)$$

Corollary 1. *For $\kappa(\mathbb{E}y, \mathbb{E}z) = \kappa(\mathbb{E}y + \mathbb{E}z)$ any optimal feasible mechanism under ex ante side contracting is a Pure Stop Loss contract.*

That is, when the loss adjustment cost only depends on the actuarial value and not the audit schedule, the insurer offers the group full marginal insurance for group losses above some aggregate group deductible. The intuition is as follows. When κ is additive in its arguments there is no need to ever audit just one agent as the additional cost of auditing the other agent is zero. Program 5 therefore reduces mathematically to that considered by Arrow (1963) where our representative group agent with preferences represented by u_0 takes the place of Arrow's policyholder.

3.1 A comparison with bilateral insurance

Instead of offering a multilateral contract to both agents, an insurance company could arrange for a bilateral insurance contract to be signed with each agent. (In a bilateral mechanism the net transfer from each agent θ_i to the insurer must depend only on that agent's loss x_i .) How does our optimal contract differ from a set of bilateral contracts? Any set of bilateral contracts offering full marginal insurance to each agent below an agent-specific deductible D_i may be written as a Generalised Stop Loss contract with $D_1 + D_2 = D_{12}$. That is, the aggregate transfer from agents to insurer satisfies $\theta_0(x) = p_0 - \max(0, x_1 - D_1) - \max(0, x_2 - D_2)$ for some deductibles $D_1, D_2 \in [0, \bar{x}]$. Such a bilateral contract is weakly dominated by the optimal multilateral contract, with strict dominance when in an optimal multilateral mechanism, $D_{12} \neq D_1 + D_2$.

A suite of bilateral contracts (left panel of Figure 1) is strictly dominated by a multilateral contract if either the deadweight loss adjustment cost of increasing transfers in the double audit region is high (right panel of Figure 1) or low (central panel of Figure 1). In practice, we might expect the latter to be more likely, as the cost of auditing a second loss from the small community may not be much higher than the cost of auditing just the first loss.

4 VERIFIABLE SABOTAGE AND SAMPLE-BASED INDEX INSURANCE

Suppose that during an audit the insurer could separately identify sabotage decisions s_1 and losses x_1 for agent 1, and condition transfer function θ on this information, but such separate identification was not possible for agent 2's losses. The indemnity schedule would no longer need to be restricted to feature no marginal overinsurance for x_1 ; the insurer could ensure that s_1 was always zero by setting θ_0 to be p_0 whenever s_1 was observed to be nonzero. However by

the following lemma, the indemnity would still need to be restricted to feature no marginal overinsurance for loss x_2 .

$$x_0 + \theta_0(x) \text{ is weakly increasing in } x_2 \quad (10)$$

Lemma 4. *Any direct feasible mechanism (a, θ) under ex ante side contracting where the insurer can separately identify sabotage decisions and losses for agent 1, but not for agent 2, is weakly dominated by a direct feasible mechanism (a', θ') that satisfies condition (10)*

Maximisation Program 5 would become:

$$\max_{p_0, y_1, y_2, z} \mathbb{E}u_0(c_0) \text{ subject to } \pi = \lambda, (4c) \text{ and } (10) \quad (11)$$

and the following Theorem characterises the form of optimal mechanisms:

Theorem 2. *In any optimal feasible mechanism under ex ante side contracting with verifiable sabotage there exist constants $D_1, D_2 \in [0, \bar{x}]$ and $D_{12} \in [0, 2\bar{x}]$ such that for almost all x :*

1. $\theta_0(x) = p_0$ for all $x_1 + x_2 \leq D_{12}$ and $x_i \leq D_i, i = 1, 2$;
2. $x_1 - y_1(x_1)$ is nonincreasing in x_1 for all $x_1 > D_1$;
3. $y_2(x_2) = \max(x_2 - D_2, 0)$;
4. $z(x) = \max[0, x_1 + x_2 - \max(y_1(x_1), y_2(x_2)) - D_{12}]$;

The optimal contract therefore retains some features from the Generalised Stop Loss contract: the equality of consumption on the double audit region; the full marginal indemnification for agent 2's losses above a single loss deductible of D_2 ; and the shape of the zero audit region. However, marginal overinsurance is offered for the first loss, above a first loss deductible of D_1 . This marginal overinsurance is for two reasons. First, losses are affiliated and so an increase in loss x_1 also implies that loss x_2 is likely to be greater. As an extreme example,

suppose that losses are close to being perfectly correlated and $D_2\kappa(Y, Z) > D_1\kappa(Y, Z)$ for all $Y, Z > 0$. Then the double audit region will be the null set and on the single audit region for agent 1, $\partial y_1(x_1)/\partial x_1 \approx 2 > 1$. The second reason for marginal overinsurance is that, since y_2 is increasing, as y_1 increases the set $\{x_2 | y_1 \geq y_2(x_2)\}$ expands in the direction of states with larger losses. This provides further impetus for an additional increase in y_1 .

In the case in which loss affiliation is strong this optimal contract may be likened to model plot area yield index agricultural insurance. The insurer would only ever conduct audits on model plots, and the aggregate transfer to agents would be a multiple of the average loss incurred on local model plots above the single loss deductible. The insurance claim payment is therefore a function of a statistical sample of local plots.

The optimal contract of Theorem 2 would only be suitable for an insurance company to offer if it could guarantee that plots could not be sabotaged without the knowledge of the insurer. If sabotage were possible, in some states of the world agents would want to create further damage, insofar as the insurer could not differentiate between the initial loss and the extra damage.

5 INDICES AND STOP LOSS GAP INSURANCE

Suppose now that there is a index $v \in V = [0, \bar{v}]$ which is jointly affiliated with the losses and costless for the insurer and agents to observe. A state of the world is now a triplet $\omega = (x_1, x_2, v) \in \Omega = X \times V$ with probability density function $f(\omega)$ and the joint affiliation assumption may be written as:

$$f(\omega)f(\omega') \leq f(\omega \vee \omega')f(\omega \wedge \omega') \text{ for all } \omega, \omega' \in \Omega \quad (12)$$

A direct mechanism will be a pair (a, θ) as before, but where the audit rule and claims transfer function can also depend on the index.

We could again reduce attention to mechanisms which satisfy the no marginal overinsurance of incurred losses condition (3) and the following revised equations (4a) and (4c):

$$\theta_0(\omega) = p_0 - I(v) - \max(y_1(x_1, v), y_2(x_2, v)) - z(x, v) \quad (13a)$$

$$p_0, I(v), y_1(x_1, v), y_2(x_2, v), z(x, v) \geq 0 \text{ for all } \omega \in \Omega \quad (13b)$$

A natural extension to assumption 1 would be that the deadweight loss adjustment cost to the insurer would be denoted by $\kappa(\mathbb{E}I, \mathbb{E}y, \mathbb{E}z)$ where $\kappa(0, 0, 0) \geq 0$ and $D_i \kappa(I, Y, Z)$ is weakly increasing in $i = 1, 2, 3$ for all $I, Y, Z \geq 0$.

Doherty and Richter (2002) considered a similar model when there was only one loss and the friction was ex ante moral hazard, rather than costly loss adjustment. They considered insurance products which offered a claim payout that was a linear function of the index plus a linear function of the *gap*, which they defined to be the incurred loss minus the indexed payout. The optimal contracts in our setting will be the multidimensional extension of Doherty and Richter's (2002) index plus gap insurance.

Following the terminology of Doherty and Richter (2002) a mechanism (a, θ) is said to offer Index Plus Generalised Stop Loss Gap insurance if:

$$\theta_0(\omega) = p_0 - \max(I(v), x_1 - D_1, x_2 - D_2, x_0 - D_{12}) \text{ for almost all } \omega \in \Omega \quad (14)$$

for some $I(v) : V \rightarrow [0, \infty)$, $D_1, D_2 \in [0, \bar{x}]$ and $D_{12} \in [0, 2\bar{x}]$. Such a composite contract offers an indexed payout of $I(v)$ but if individual or joint losses are large enough there will be an indemnity based top-up of $\max(0, x_1 - I(v) - D_1, x_2 - I(v) - D_2, x_0 - I(v) - D_{12})$. Net consumption of the pair of agents for almost all $\omega \in \Omega$ is:

$$c_0(\omega) = w_0 - p_0 - \min(x_1 + x_2 - I(v), D_1 + x_2, x_1 + D_2, D_{12}) \quad (15)$$

Theorem 3. *Any optimal feasible mechanism under ex ante side contracting with a costlessly observable index offers Index Plus Generalised Stop Loss Gap insurance. Any optimal index claim function is weakly increasing in the index realisation.*

The general form of the optimal contract optimal contract in this case is an extension to the Generalised Stop Loss contract of section 3, but where indemnity payments are based on total audited losses net of the index payment $I(v)$.

A mechanism (a, θ) is said to offer Index Plus Pure Stop Loss Gap insurance if:

$$\theta_0(\omega) = p_0 - \max(I(v), x_0 - D_{12}) \text{ for almost all } \omega \in \Omega \quad (16)$$

Corollary 2. *For $\kappa(\mathbb{E}I, \mathbb{E}y, \mathbb{E}z) = \kappa(\mathbb{E}I, \mathbb{E}y + \mathbb{E}z)$ any optimal feasible mechanism under ex ante side contracting with a costlessly observable index offers Index Plus Pure Stop Loss Gap insurance*

That is to say, in the optimal contract the group receives an indexed payment $I(v)$ which depends only on the realised index and indemnity based payouts provide a floor of D_{12} on aggregate net loss $x_0 - I(v)$.

6 CROWDING IN

Although there is ample evidence that the poor find affordable ways to share risk within close knit groups, such as households, extended families or villages, such pooling is rarely observed to be Pareto optimal.¹⁰ A side contract where agents can only commit to the transfer rule after the state of the world x is realised will

¹⁰ There is ample evidence that mutual insurance within close knit groups, such as households, extended families or villages departs from first best (e.g. Townsend 1994, Udry 1994, Dercon 2002, Fafchamps and Lund 2003). Informal punishments such as exclusion from nonmarket insurance, social sanctions or physical violence may be able to induce small, but not large, insurance transfers (Coate and Ravallion 1993). Empirical investigations such as Ligon, Thomas and Worrall (2002) provide evidence for the existence of binding enforcement constraints in poor communities, causing mutual insurance arrangements to depart from first best in the direction of Coate and Ravallion (1993).

be called an *interim side contract* $S_x = (\{s, m, \tau\}_{i=1,2})$ where $s \in X$, $m_i \in M_i$ and $\tau_i \in \mathbb{R}$. Abusing notation, a collection of interim side contracts $\{S_x\}_{x \in X}$ will be denoted S . In the absence of incentives from the insurer, interim side contracting does not allow losses to be pooled between group members; by the interim stage the loss uncertainty has already been resolved and a lucky agent has no incentive to agree to make transfers to an unlucky agent. However, agents will still be able to collude against the insurer, by sending messages that maximise the total consumption and undoing any contractual split of consumption dictated by the split of θ_i .

Following Rai and Sjostrom (2004) we will show that any outcome that can be implemented under ex-ante side contracting can also be implemented under interim side contracting, so long as the insurer has the ability to inflict a large enough non-pecuniary punishment $\bar{q}$ on each agent. Let the non-pecuniary punishment for agent i be $q_i : M_1 \times M_2 \times X \rightarrow [0, \bar{q}]$. This could be the denial of future financial services or the cost of being ‘hassled’ by the insurer. Each agent i is assumed to have preferences defined over net consumption c_i and q_i with $u_i(c_i, q_i) = u_i(c_i - q_i)$. A mechanism under interim side contracting is taken to be a specification of (M, a, θ, q) .

For a given mechanism G , an interim side contract S will be called incentive compatible if the agents could not sign a side contract that would be better for both agents in any state of the world. Individual rationality requires each side contract S_x to be accepted by both agents for each state $x \in X$. If the agents do not sign any side contract at the interim stage then they go on to play a Nash equilibrium of (G, x) . Let $u_i(G, x)$ denote agent i ’s Nash equilibrium payoff.¹¹ $u_i(G, x)$ acts as the reservation utility for interim side contracting in state x , and so we may write the individual rationality and incentive compatibility constraints as:

¹¹ If there were multiple Nash equilibria of (G, x) we would assume that agents make some selection from the set of Pareto optimal Nash equilibria. However, the mechanism we construct has a unique Nash equilibrium.

(IR2) (For individual rationality under interim side contracting)

$$\mathbb{E}[\theta_0] - \kappa \geq 0, \mathbb{E}[u_i(c_i)] \geq \bar{u}_i \text{ and } u_i(c_i(x)) \geq u_i(G, x) \text{ for } i = 1, 2 \text{ and } x \in X.$$

(IC3) (For incentive compatibility of an interim side contract)

There is no other individually rational side contract which gives strictly higher utilities to both agents in any state of the world.

A mechanism G will be called *feasible under interim side contracting* if there exists some ex ante side contract S which, together with G , satisfies (IR2), (IC2) and (IC3).

Assuming that the maximum punishment $\bar{q}$ is at least as large as the largest side transfer of the outcome we intent to implement, we may show the following:

Theorem 4. *Any outcome (π^*, c_1^*, c_2^*) from an optimal feasible mechanism under ex ante side contracting, and associated feasible ex ante side contract, is implementable by a feasible mechanism and interim side contract. Moreover, it is uniquely ϵ -implementable in the following sense: for any $\epsilon > 0$ there exists a mechanism $G(\epsilon)$ such that for any individually rational and incentive compatible side contract S , $|c_1^*(x) - c_1(x)|, |c_2^*(x) - c_2(x)| \leq \epsilon$.*

This result is equivalent to Proposition 3 of Rai and Sjostrom (2004) and relies on both the ability of the insurer to punish and that of agents to cross report. We show that Generalised Stop Loss contracts are still attractive even if agents cannot commit ex ante to pool uncertainty, so long as agents have a loss adjustment cost advantage and the insurer has sufficient ability to punish agents.

7 CONCLUSION

How is microinsurance, that is insurance for low-income people, economically different to conventional personal lines insurance? The key assumption modeled

in this paper is that for microinsurance the cost of ex-post claims processing, known as *loss adjustment*, is lower for local nonmarket institutions relative to external formal sector insurers. Although multilateral credit contracts are now common, multilateral insurance contracts, where the claim payment to one policyholder explicitly depends on the losses incurred by other policyholders, are rarely observed in practice, with the exception of area yield agricultural insurance (Mahul and Stutley 2010). A normative interpretation of the above models provides support for particular multilateral insurance contract forms between premium-charging insurance companies or publicly funded social insurance schemes and poor individuals, namely the Stop Loss, Sample-Based Index and Index Plus Gap contract forms.

Development economists have a healthy suspicion of normative microeconomic theory brought about by the observation that the poor are usually more enterprising than the researcher. However, this suspicion is less well-founded in the context of formal and semiformal finance for the poor. The potential welfare gains from successful financial innovation are widely considered to be large (Banerjee 2002, Collins et al. 2009, Karlan and Morduch 2009). However, experimenting with financial innovations is costly and outcomes are difficult to evaluate, particularly for insurance contracts where payouts are typically expected to be made in only one year out of five. Financial innovations in developed markets have often been theory-led and the recent advances in positive microfinance theory leave economists well placed to make suggestions for improvements.

In their chapter on microfinance in the Handbook of Development Economics, Karlan and Morduch (2009) write: “[Micro-insurance] holds promise, but the field is young and no approaches have emerged so far that offer break-throughs akin to the original group-lending innovations that ignited the global explosion of microcredit”. We wonder whether the contracts outlined in this paper might be useful here and there for formal contracting with the poor, particularly in the life, longevity and crop insurance classes of business.

REFERENCES

- Arnott, R. and J.E. Stiglitz**, “Moral Hazard and Nonmarket Institutions: Dysfunctional Crowding Out of Peer Monitoring?,” *The American Economic Review*, 1991, 81 (1), 179–190.
- Arrow, Kenneth J.**, “Uncertainty and the Welfare Economics of Medical Care,” *The American Economic Review*, 1963, 53 (5), 941–973.
- Banerjee, A.**, “The Uses of Economic Theory: Against a Purely Positive Interpretation of Theoretical Results,” Working Paper, Massachusetts Institute of Technology, Dept. of Economics 2002.
- Behrman, J.R. and A.B. Deolalikar**, “Health and nutrition,” *Handbook of development economics*, 1988, 1, 631–711.
- **and** ———, “The intrahousehold demand for nutrients in rural south India: Individual estimates, fixed effects, and permanent income,” *Journal of human resources*, 1990, 25 (4), 665–696.
- Bennett, P.**, “Mutuality at a distance? Risk and regulation in marine insurance clubs,” *Environment and Planning A*, 2000, 32 (1), 147–164.
- Bond, E.W. and K.J. Crocker**, “Hardball and the soft touch: The economics of optimal insurance contracts with costly state verification and endogenous monitoring costs,” *Journal of Public Economics*, 1997, 63 (2), 239–264.
- Borch, Karl**, “Equilibrium in a Reinsurance Market,” *Econometrica*, 1962, 30 (3), 424–444.
- Cabrales, A., A. Calvo-Armengol, and M.O. Jackson**, “La crema: A case study of mutual fire insurance,” *Journal of Political Economy*, 2003, 111 (2), 425–458.
- Clarke, D.J.**, “A Theory of Rational Demand for Index Insurance,” Department of Economics Discussion Paper Series 572, University of Oxford 2011.
- Coate, S. and M. Ravallion**, “Reciprocity Without Commitment,” *Journal of Development Economics*, 1993, 40 (1), 1–24.
- Collins, D., J. Morduch, S. Rutherford, and O. Ruthven**, *Portfolios of the Poor: How the World’s Poor Live on \$2 a Day*, Princeton University Press, 2009.
- Crocker, K.J. and J. Morgan**, “Is honesty the best policy? Curtailing insurance fraud through optimal incentive contracts,” *Journal of Political Economy*, 1998, 106 (2), 355–375.
- Dercon, S.**, “Income Risk, Coping Strategies, and Safety Nets,” *The World Bank Research Observer*, 2002, 17 (2), 141–166.
- , “Risk, insurance, and poverty: a review,” in S. Dercon, ed., *Insurance against Poverty*, 2004, chapter 1.
- **and P. Krishnan**, “In Sickness and in Health: Risk Sharing within Households in Rural Ethiopia,” *The Journal of Political Economy*, 2000, 108 (4), 688–727.
- , **J. De Weerd, T. Bold, and A. Pankhurst**, “Group-based funeral insurance in Ethiopia and Tanzania,” *World Development*, 2006, 34 (4), 685–703.

- Derrig, R.A.**, “Insurance fraud,” *Journal of Risk and Insurance*, 2002, pp. 271–287.
- Dionne, G., ed.**, *Handbook of Insurance*, Kluwer Academic, 2000.
- Doherty, N.A. and A. Richter**, “Moral hazard, basis risk, and gap insurance,” *Journal of Risk and Insurance*, 2002, pp. 9–24.
- Fafchamps, M. and S. Lund**, “Risk-sharing networks in rural Philippines,” *Journal of Development Economics*, 2003, 71, 261–287.
- Fagart, M.C. and P. Picard**, “Optimal insurance under random auditing,” *The GENEVA Papers on Risk and Insurance-Theory*, 1999, 24 (1), 29–54.
- Ghatak, M. and TW Guinnane**, “The economics of lending with joint liability: theory and practice,” *Journal of Development Economics*, 1999, 60 (1), 195–228.
- Gollier, C. and H. Schlesinger**, “Arrow’s Theorem on the optimality of deductibles: A stochastic dominance approach,” *Economic Theory*, 1996, 7 (2), 359–363.
- Hölmstrom, B.**, “Moral Hazard and Observability,” *Bell Journal of Economics*, 1979, 10 (1), 74–91.
- Ibarra, H.**, “Self-Insurance Funds in Mexico,” in E.N. Gurenko, ed., *Catastrophe Risk and Reinsurance: A Country Risk Management Perspective*, 2004, chapter 18.
- Jacoby, H.G. and E. Skoufias**, “Risk, financial markets, and human capital in a developing country,” *The Review of Economic Studies*, 1997, pp. 311–335.
- Jütting, J.P.**, “Do community-based health insurance schemes improve poor people’s access to health care? Evidence from rural Senegal,” *World Development*, 2004, 32 (2), 273–288.
- Karlan, D. and J. Morduch**, “Access to Finance: Credit Markets, Insurance, and Saving,” in Dani Rodrik and Mark Rosenzweig, eds., *Handbook of Development Economics*, Vol. 5, Amsterdam: North Holland, 2009, chapter 71.
- Krasa, S. and A.P. Villamil**, “Optimal contracts when enforcement is a decision variable,” *Econometrica*, 2000, pp. 119–134.
- Laffont, J.J. and D. Martimort**, “Mechanism Design with Collusion and Correlation,” *Econometrica*, 2000, 68 (2), 309–342.
- Ligon, E., J.P. Thomas, and T. Worrall**, “Informal insurance arrangements with limited commitment: Theory and evidence from village economies,” *Review of Economic Studies*, 2002, 69 (1), 209–244.
- Mahul, O.**, “Optimum area yield crop insurance,” *American Journal of Agricultural Economics*, 1999, 81 (1), 75.
- **and C.J. Stutley**, *Government Support to Agricultural Insurance: Challenges and Options for Developing Countries*, World Bank Publications, 2010.
- Milgrom, P.R. and R.J. Weber**, “A Theory of Auctions and Competitive Bidding,” *Econometrica*, 1982, 50 (5), 1089–1122.

- Miranda, M.J.**, “Area-yield crop insurance reconsidered,” *American Journal of Agricultural Economics*, 1991, 73 (2), 233–242.
- Morduch, J.**, “Between the State and the Market: Can Informal Insurance Patch the Safety Net?,” *The World Bank Research Observer*, 2002, 14 (2), 187–207.
- Picard, P.**, “On the design of optimal insurance policies under manipulation of audit cost,” *International Economic Review*, 2000, pp. 1049–1071.
- Rai, A.S. and T. Sjöstrom**, “Is Grameen Lending Efficient? Repayment Incentives and Insurance in Village Economies,” *Review of Economic Studies*, 2004, 71 (1), 217–234.
- Rothschild, Michael and Joseph E. Stiglitz**, “Increasing risk: I. A definition,” *Journal of Economic Theory*, September 1970, 2 (3), 225–243.
- Swiss Re**, “An Introduction to Reinsurance,” Technical Report, Swiss Re Technical Publishing, Zurich 2002.
- Townsend, Robert M.**, “Optimal Contracts and Competitive Markets with Costly State Verification,” *Journal of Economic Theory*, 1979, 21 (2), 265–293.
- , “Risk and Insurance in Village India,” *Econometrica*, 1994, 62 (3), 539–591.
- Udry, C.**, “Risk and Insurance in a Rural Credit Market: An Empirical Investigation in Northern Nigeria,” *The Review of Economic Studies*, 1994, 61 (3), 495–526.
- Wilson, R.**, “The Theory of Syndicates,” *Econometrica*, 1968, 36 (1), 119–132.
- Winton, A.**, “Costly state verification and multiple investors: the role of seniority,” *Review of Financial Studies*, 1995, 8 (1), 91–123.

APPENDIX

Proof of Proposition 1. Consider the following direct mechanism. Let each agent report the total loss $s + x$, so $M_1 = M_2 = X$. Define $\theta_i^*(x) = w_i - x_i - c_i(x)$ where $(c_1(x), c_2(x))$ is the desired consumption schedule. If both agents report the same total loss x then neither agent is audited ($a_1(x, x) = a_2(x, x) = 0$) and transfers are such each agent's consumption follows the desired consumption schedule ($\theta_i(x, x, x') = \theta_i^*(x)$).

If agents report different states, then both agents are audited ($a_1(x', x'') = a_2(x', x'') = 1$ where $x' \neq x''$) and the insurer learns which agents have lied. If both agents have lied ($x' \neq x + s$ and $x'' \neq x + s$) then set $\theta_i(x', x'', x + s) = \theta_i^*(x + s) + \epsilon$ for some $\epsilon > 0$. If agent 1 tells the truth and agent 2 doesn't then set $\theta_1(x, x'', x + s) = \min(\theta_1^*(x''), \theta_1^*(x + s)) - \epsilon$ and $\theta_2(x + s, x'', x + s) = \theta_2^*(x + s) + \epsilon$. If agent 2 tells the truth and agent 1 doesn't then set $\theta_2(x', x + s, x + s) = \min(\theta_2^*(x'), \theta_2^*(x + s)) - \epsilon$ and $\theta_1(x', x + s, x + s) = \theta_1^*(x + s) + \epsilon$.

Under this mechanism misreporting the total loss ($m_i \neq x + s$) is strictly dominated by truthtelling ($m_i = x + s$) for both agents. If $c_i(x)$ is nonincreasing in x_i for $i = 1, 2$ then there is a Nash equilibrium where both agents choose $s_i = 0$ and then report truthfully, with no auditing on the equilibrium path. However, there are equilibria with sabotage by agent i in regions where $c_i(x)$ is constant over x_i . If $c_i(x)$ is strictly decreasing in x_i for $i = 1, 2$ then the no sabotage, truthtelling equilibrium is the unique equilibrium. $\square$

Proof of Lemma 1. Define $\sigma : X \rightarrow X$ by $\sigma(x) = (x_1 + s_1(x), x_2 + s_2(x)) \forall x \in X$, $a' = a \circ (\sigma, \sigma)$, $\theta' = \theta \circ (\sigma, \sigma, (\sigma, Id))$, where θ is now a function of $m_1, m_2, (x, s_1)$.

Now for any feasible ex ante side contract $S = (s, \tau)$ we define $S' = (s', \tau)$ where $s'(x) = (0, 0)$ for all $x \in X$. Direct mechanism (a', θ') inherits feasibility from (a, θ) and under side contract S' both agents receive the same expected utility as under the original mechanism and S , but the insurer receives weakly higher expected profits as net transfers to agents are lower in states where $s(x) \neq (0, 0)$ under the original side contract. $\square$

Proof of Lemma 2. For feasible (a, θ_0) define

$$p_0 := \max_{x \in X} \theta_0(x) \tag{A-1a}$$

$$y_i(x_i) := \min_{x_j \in [0, \bar{x}]} (p_0 - \theta_0(x)) , (i, j) \in \{(1, 2), (2, 1)\} \tag{A-1b}$$

$$z(x) := p_0 - \max(y_1(x_1), y_2(x_2)) - \theta_0(x) \tag{A-1c}$$

(A-1c) ensures that definitions (A-1a)-(A-1c) satisfy (4a).

To demonstrate that they satisfy (4c), first suppose $p_0 < 0$. Then $\theta_0(x) < 0$ for all $x \in X$, violating (IR1). So $p_0 \geq 0$. Since $\max_{x \in X} \theta_0(x) - \theta_0(x) \geq 0$, $y_i(x_i) \geq 0$ for $i = 1, 2$ and $x \in X$. Finally, substituting (A-1b) in to (A-1c) gives

$$\begin{aligned} z(x) &= -\max \left(\min_{x'_1 \in [0, \bar{x}]} (-\theta_0(x'_1, x_2)), \min_{x'_2 \in [0, \bar{x}]} (-\theta_0(x_1, x'_2)) \right) - \theta_0(x) \\ &= \min \left(\max_{x'_1 \in [0, \bar{x}]} (\theta_0(x'_1, x_2)), \max_{x'_2 \in [0, \bar{x}]} (\theta_0(x_1, x'_2)) \right) - \theta_0(x) \\ &\geq 0 \end{aligned}$$

where the final inequality arises from $\max_{x'_1 \in [0, \bar{x}]} (\theta_0(x'_1, x_2)) \geq \theta_0(x)$ and $\max_{x'_2 \in [0, \bar{x}]} (\theta_0(x_1, x'_2)) \geq \theta_0(x)$. Definitions (A-1a)-(A-1c) therefore satisfy (4c).

Now we show by contradiction that feasibility implies that 4b must hold for (A-1a)-(A-1c). First suppose that there is some $y_1(x_1) > y_2(x_2)$ such that $a_1(x_1, x_2) = 0$. $y_1(x_1) \leq p_0 - \theta_0(x_1, x_2)$ from (A-1b), with strict inequality if there was some $x'_2 \in [0, \bar{x}]$ such that $\theta_0(x_1, x'_2) > \theta_0(x_1, x_2)$. For (IC2) to hold $\theta_0(x'_1, x_2) \leq \theta_0(x_1, x_2)$ for all $x'_1 \in [0, \bar{x}]$ otherwise it would be optimal for agents to misreport state (x'_1, x_2) as (x_1, x_2) . Now, from (A-1b), $y_2(x_2) := \min_{x'_1 \in [0, \bar{x}]} (p_0 - \theta_0(x'_1, x_2)) = p_0 - \theta_0(x_1, x_2)$. However, combined with $y_1(x_1) \leq p_0 - \theta_0(x_1, x_2)$ it must be that $y_1(x_1) \leq y_2(x_2)$. Contradiction. The case with some $y_1(x_1) > y_2(x_2)$ such that $a_1(x_1, x_2) = 0$ follows trivially.

Second, suppose that there is some $x \in X$ such that $y_1(x_1) = y_2(x_2) > 0$ but $a_1(x) = a_2(x) = 0$. For (IC2) to hold $\theta_0(x') \leq \theta_0(x)$ for all $x' \in X$ otherwise it would be optimal for agents to misreport state x' as x . Therefore $p_0 = \theta_0(x)$ and $y_1(x_1) = y_2(x_2) = 0$ from (A-1a) and (A-1b). Contradiction.

Third, suppose that there is some $x \in X$ such that $z(x) > 0$ but $a_1(x) = 0$. For (IC2) to hold $\theta_0(x'_1, x_2) \leq \theta_0(x_1, x_2)$ for all $x'_1 \in [0, \bar{x}]$ otherwise it would be optimal for agents to misreport state (x'_1, x_2) as (x_1, x_2) . Now, from (A-1b), $y_2(x_2) := \min_{x'_1 \in [0, \bar{x}]} (p_0 - \theta_0(x'_1, x_2)) = p_0 - \theta_0(x_1, x_2)$ and so $z(x) = 0$ from (A-1c). Contradiction. The case with some $x \in X$ such that $z(x) > 0$ and $a_2(x) = 0$ follows trivially. $\square$

Proof of Lemma 3. We may parameterise any constrained Pareto optimal mechanism and side contract by the expected utility of agent 2, μ , and the expected profit of the insurer, λ :

$$\begin{aligned} & \max_{p_0, y_1, y_2, z, c_1, c_2} \mathbb{E}u_1(c_1(x)) \\ & \text{subject to} \\ & \mathbb{E}u_2(c_2(x)) \geq \mu \\ & \pi(p_0, \mathbb{E}y(x), \mathbb{E}z(x)) \geq \lambda \\ & \text{(IC1), (IC2), } y = \max(y_1, y_2), \theta_0 = p_0 - y - z \\ & \text{and } c_1 + c_2 = w_0 - p_0 - x_0 + y + z \end{aligned}$$

We may replace (IC1) and (IC2) with (3) and (4c). Both of these depend on aggregate consumption, but not the split of consumption between agents.

Denoting the Lagrangian multiplier on the budget constraint as $f(x) \times u'_0(c_0(x))$ for some function u'_0 , and that on the expected utility constraint for agent 2 as ν , the first order constraints for c_1 and c_2 yield:

$$u'_1(c_1^*(x)) = \nu u'_2(c_0^*(x) - c_1^*(x)) = u'_0(c_0^*(x)) \quad \forall x \in X$$

We may integrate u'_0 to construct a function u_0 such that $u_0(c_0^*(x)) = u_1(c_1^*(x))$ for all x . u_0 is increasing and strictly concave from the strict increasing concavity of u_1 and u_2 . The required result follows by substituting the expression for $u_0(c_0(x))$ into the objective function. $\square$

Proof of Theorem 1. Throughout the proof we fix p_0 , $\mathbb{E}y$ and $\mathbb{E}z$, and therefore the insurer's expected profit, and appeal to notions of second order stochastic dominance

of aggregate consumption schedule c_0 .¹² Suppose $\{p_0, y_1, y_2, z\}$ solved Program 5 but was not a Generalised Stop Loss contract.

First we show that:

$$z(x) = \max[0, x_1 + x_2 - \max(y_1(x_1), y_2(x_2)) - D_{12}] \quad (\text{A-2})$$

for some D_{12} , offering a floor on consumption. Fix y_1 and y_2 and suppose that z does not satisfy equation (A-2). Define z' as in equation (A-2) for some D'_{12} chosen so that $\mathbb{E}z = \mathbb{E}z'$. Abusing notation, define random variables $c_0 = w_0 - p_0 - x_0 + y + z$ and $c'_0 = w_0 - p_0 - x_0 + y + z'$ with cdfs G_{c_0} and $G_{c'_0}$ respectively. By Gollier and Schlesinger (1996), c_0 is a mean preserving spread of c'_0 and so $G_{c'_0}$ strictly SSD G_{c_0} . c'_0 will therefore be strictly preferred by the representative group agent, whose preferences are strictly risk averse, and therefore by both agents.

To characterise optimal y_1, y_2 we appeal to the following Lemma to reduce attention to net consumption before the addition of z , $c_0^{-z} = w_0 - p_0 - x_0 + \max(y_1, y_2)$:

Lemma A1. *If c'_0 SSD c_0 and D'_{12}, D_{12} are set so that $\mathbb{E} \max(0, D'_{12} - c'_0) = \mathbb{E} \max(0, D_{12} - c_0)$ then $\max(c'_0, D'_{12})$ SSD $\max(c_0, D_{12})$. The dominance is strict iff the cdfs of $\max(c'_0, D'_{12})$ and $\max(c_0, D_{12})$ are not identical.*

Proof. $G_{c'_0}$ SSD G_{c_0} iff $\int_{-\infty}^t [G_{c_0}(c) - G_{c'_0}(c)]dc \geq 0 \forall t$. We are also given that D'_{12}, D_{12} are set so that

$$\int_{-\infty}^{D'_{12}} G_{c'_0}(c)dc = \int_{-\infty}^{D_{12}} G_{c_0}(c)dc \quad (\text{A-3})$$

Now:

$$\begin{aligned} \int_{-\infty}^t [G_{\max(c_0, D_{12})}(c) - G_{\max(c'_0, D'_{12})}(c)]dc &= \int_{D_{12}}^t G_{c_0}(c)dc - \int_{D'_{12}}^t G_{c'_0}(c)dc \\ &= \int_{-\infty}^t [G_{c_0}(c) - G_{c'_0}(c)]dc \geq 0 \forall t \end{aligned}$$

Where the second equality is from equation (A-3) □

Consider the following definitions:

$$\begin{aligned} y'(x) &:= \max(y_j(x_j), x_i - D_i(x_j)) \\ y''(x) &:= \max(y_j(x_j), y''_i(x_i)) \text{ where } y''_i(x_i) = \max(x_i - D_i, 0) \end{aligned}$$

where D_i is chosen so that $\int_X f(x)[y''(x) - y(x)]dx = 0$ and $D_i(x_j)$ is chosen so that $\int_0^{\bar{x}} f(x)[y'(x) - y(x)]dx_i = 0$ for each x_j such that $\mathbb{E}[y(x)|x_j] > 0$. Due to our assumption (1) that losses are affiliated, $D_i(x_j)$ must be weakly increasing in x_j for those x_j such that $\mathbb{E}[y(x)|x_j] > 0$. We may therefore define $D_i(x_j)$ for those x_j for which $\mathbb{E}[y(x)|x_j] = 0$ to be no greater than $D_i(x'_j)$ for all $x'_j > x_j$ and no less than $D_i(x''_j)$ for all $x'_j < x_j$.

¹² Following Rothschild and Stiglitz (1970), second order stochastic dominance is defined as follows:

$$H \text{ SSD } G \text{ iff } \int_{-\infty}^t [G(c) - H(c)]dc \geq 0 \forall t$$

with strict SSD if there is a t for which the inequality is strict.

Let $c'_0(x) := w_0 - p_0 - x_0 + y'(x)$ and $c''_0(x) := w_0 - p_0 - x_0 + y''(x)$. We will show that $G_{c'_0} \text{ SSD } G_{c_0} \text{ strict SSD } G_{c_0^-}$ and that the new contract $\{p_0, y'_i, y_j, z\}$ is feasible and yields the same expected profit for the insurer as the original contract. This will provide us with our contradiction.

$\{p_0, y'_i, y_j, z\}$ is feasible and yields the same expected profit for the insurer by construction so all we need prove are the SSD relationships. That $G_{c'_0} \text{ strict SSD } G_{c_0^-}$ is straightforward: we perform a series of mean preserving contractions of consumption, one for each x_j , by moving claims mass to equalise net consumption in the highest net loss states. The strictness arises from the observation that if $G_{c'_0}$ equals $G_{c_0^-}$ then the original contract must have been a Generalised Stop Loss contract.

Finally we show that $G_{c'_0} \text{ SSD } G_{c_0}$. Since $D_i(x_j)$ is weakly decreasing in x_j there is some x_j^* such that $D_i(x_j) \leq D_i$ for all $x_j \leq x_j^*$ and $D_i(x_j) \geq D_i$ for all $x_j \geq x_j^*$. To transform $c'_0(x)$ to $c_0(x)$ we reduce claim payments for $x_j \leq x_j^*$ and increase claim payments for $x_j \geq x_j^*$. The minimum consumption in the region $x_j \leq x_j^*$ after reducing claim payments is $w_0 - p_0 - \min(D_i + x_j^*, \bar{x} + x_j^* - y_j(x_j))$. Moreover, the maximum consumption in the region $x_j \geq x_j^*$ after reducing claim payments is also equal to $w_0 - p_0 - \min(D_i + x_j^*, \bar{x} + x_j^* - y_j(x_j))$. We have moved claims mass from high net consumption states to low net consumption states and therefore $G_{c'_0} \text{ SSD } G_{c_0}$. $\square$

Proof of Corollary 1. When κ is additive in its arguments, auditing two agents and making a claim payment costs the same as auditing only one and making the same claim payment and so without loss of generality we may assume that $y_1(x_1) = y_2(x_2) = 0$ for all $x \in X$ and instead write the optimal net transfer in terms of p_0 and $z(x) = p_0 - \theta_0(x)$ only. Theorem 1 holds, implying the required result. $\square$

Proof of Lemma 4. Define $\sigma : X \rightarrow X$ by $\sigma(x) = (x_1, x_2 + s_2(x)) \forall x \in X$, $a' = a \circ (\sigma, \sigma)$, $\theta' = \theta \circ (\sigma, \sigma)$.

Now for any feasible ex ante side contract $S = (s, \tau)$ we define $S' = (s', \tau)$ where $s'(x) = (s_1(x), 0)$ for all $x \in X$. Direct mechanism (a', θ') inherits feasibility from (a, θ) and under side contract S' both agents receive the same expected utility as under the original mechanism and S , but the insurer receives weakly higher expected profits as net transfers to agents are lower in states where $s_2(x) \neq 0$ under the original side contract. $\square$

Proof of Theorem 2. Parts 3. and 4. can be shown using the same logic as in the proof of Theorem 1. Appealing to Lemma A1 we may restrict attention to consumption before the addition of z, c_0^- . Define $y(x) = \max(y_1(x_1), y_2(x_2))$ for some optimal y_1, y_2 .

First we show that the region $y_1(x_1) = 0$ is a lower interval of x_1 , and therefore that Part 1. of the Theorem holds. For each x_2 define $F_y(y|x_2)$ as the cdf of the random variable $y(x)$ conditional on the second loss. Consider the following definitions:

$$y'(x) := \sup_{y'} \{F_y(y'|x_2) \leq F_{x_1}(x_1|x_2)\} \text{ for all } x \in X$$

$$y''(x) := \max(y'_1(x_1), y_2(x_2))$$

where $y'_1(x_1)$ is the y'_1 that solves $\int_0^{\bar{x}} f(x) y'(x) dx_2 = \int_0^{\bar{x}} f(x) \max(y'_1, y_2(x_2)) dx_2$ or zero if there is no solution, for each x_1 . Operation $y \rightarrow y'$ rearranges the claims mass to states with high x_1 , keeping the conditional cdf constant for each x_2 . Operation $y' \rightarrow y''$ shifts claims mass towards states with higher x_2 whilst making y incentive compatible, keeping $\mathbb{E}[y|x_1]$ constant.

By construction $y''(x)$ satisfies incentive compatibility and $G_{w_0-p_0-x_0+y''} SSD G_{w_0-p_0-x_0+y'} SSD G_{w_0-p_0-x_0+y}$ with strict dominance if the original y was not weakly increasing in both arguments for almost all x .

We will only provide a sketch proof to part 2. Suppose that there are some intervals $\beta = [\beta_1, \beta_2]$ and $\beta' = [\beta'_1, \beta'_2]$ such that $\beta'_2 > \beta'_1 > \beta_2 > \beta_1$ and for all $x_1 \in \beta, x'_1 \in \beta'$ then $y_1(x'_1) \geq y_1(x_1) > 0$ but $x'_1 - y_1(x'_1) > x_1 - y_1(x_1)$. By affiliation of losses and part 3. of this Theorem, the conditional cdf of consumption in region $\{x|x_1 \in \beta, y_1(x_1) > y_2(x_2)\}$ strictly SSD that of consumption in region $\{x|x_1 \in \beta', y_1(x_1) > y_2(x_2)\}$ and so we can strictly increase expected utility of the agents by decreasing $y_1(x_1)$ for all $x_1 \in \beta$ by some $\epsilon > 0$ and increasing $y_1(x_1)$ for all $x_1 \in \beta'$ by some $\epsilon' > 0$ where ϵ, ϵ' are chosen so that $\mathbb{E}y$ remains constant. $\square$

Proof of Theorem 3. Proof of Theorem 3 closely follows that of Theorem 1 and so we only provide a sketch.

Conditional on function $I(v)$ and keeping $\mathbb{E}y_1, \mathbb{E}y_2, \mathbb{E}z$ constant the insurer must choose optimal y_1, y_2, z . If the insurer audits only agent i it learns the net loss $x_i - I(v)$ but does not learn x_j , and if the insurer audits both agents it learns the complete net loss $x_1 + x_2 - I(v)$. Following the same logic in the proof of Theorem 1 the double audit transfer will act to provide a floor on the net loss $x_1 + x_2 - I(v)$ and the single audit transfers will provide a floor on $x_i - I(v)$.

All that remains to be shown is that $I(v)$ is increasing in v . We may extend Lemma A1 to show that, conditional on y_1, y_2, z being of the form above, if $G_{w_0-p_0-x_0+I'(v)} SSD G_{w_0-p_0-x_0+I(v)}$ then no mechanism with indexed payout function I can be optimal. Since v is jointly affiliated with x , any optimal indemnity schedule must be increasing in v for a.e. v or we can rearrange I to give a second order stochastically dominant consumption schedule for agents. $\square$

Proof of Corollary 2. Proof follows that of Corollary 1. $\square$

Proof of Theorem 4. Suppose the optimal direct mechanism is (a^*, θ^*) and associated ex ante side contract is (s^*, τ^*) . There is no overinsurance (Lemma 1 holds) and so $s^*(x) = 0$ for all x .

Consider the following mechanism (M, a, θ, q) and a collection of feasible interim side contracts (s, m, τ) . Under the mechanism each agent i must send a message (x^i, r^i) where x^i is i 's report of the state of the world, $r^i = 1$ is a report that the other agent j has agreed to make the transfer $\tau_j^*(x)$ to agent i , and $r^i = 0$ is a report that agent j hasn't. The message space, audit rule, net transfer to insurer and punishment function

are defined as follows for some small positive $\epsilon > 0$:

$$\begin{aligned}
 M_i &:= X \times \{0, 1\} \\
 a_i((x^1, r^1), (x^2, r^2)) &:= \begin{cases} a^*(x^i) & \text{if } x^1 = x^2 \text{ and } r^1 = r^2 = 1 \\ (1, 1) & \text{otherwise} \end{cases} \\
 \tilde{x}^i &:= \begin{cases} (x_1^i, x_2^i) & \text{if } a_1 \times (x_1^i - x_1) = 0, a_2 \times (x_2^i - x_2) = 0 \\ (x_1^i, x_2) & \text{if } a_1 \times (x_1^i - x_1) = 0, a_2 \times (x_2^i - x_2) \neq 0 \\ (x_1, x_2^i) & \text{if } a_1 \times (x_1^i - x_1) \neq 0, a_2 \times (x_2^i - x_2) = 0 \\ (x_1, x_2) & \text{if } a_1 \times (x_1^i - x_1) \neq 0, a_2 \times (x_2^i - x_2) \neq 0 \end{cases} \\
 \theta_i((x^1, r^1), (x^2, r^2), x) &:= \begin{cases} \theta_i^*(\tilde{x}^i) & \text{if } r^i = 1 \\ \theta_i^*(\tilde{x}^i) + \min(0, \tau_i^*(\tilde{x}^i)) & \text{if } r^i = 0 \end{cases} \\
 q_i((x^1, r^1), (x^2, r^2), x) &:= \begin{cases} \bar{q} & \text{if } \tilde{x}^i \neq x^i \\ \max(0, \tau_i^*(\tilde{x}^i)) + \epsilon & \text{if } \tilde{x}^i = x^i \text{ and } r^j \times r^i = 0 \\ 0 & \text{if } \tilde{x}^i = x^i \text{ and } r^j = r^i = 1 \end{cases}
 \end{aligned}$$

So we define $\tilde{x}^i$ as the state of the world reported by agent i , x^i , corrected for any information discovered through auditing.

For a particular state x , consider the reservation utility of agent i . No matter what the message and transfer of agent j , agent i can receive utility of $u_i(c_i^*(x) - \epsilon)$, where $c_i^*(x) = w_i - x_i - \theta_i^*(x) - \tau_i^*(x)$. If $\tau_i^*(x) > 0$ then agent i can send message $m_i = (x, 1)$ to the insurer and make zero transfer to agent j . Agent i would transfer $\theta_i^*(x)$ to the insurer and receive a maximum punishment of $\tau_i^*(x) + \epsilon$. If $\tau_i^*(x) \leq 0$ then agent i can send message $m_i = (x, 0)$ to the insurer and make zero transfer to agent j . Agent i would transfer $\theta_i^*(x) + \tau_i^*(x)$ to the insurer and receive a maximum punishment of ϵ . In either case the utility of the agent is at least $u_i(c_i^*(x) - \epsilon)$.

Any interim side contract satisfying (IR2) must yield utility for each agent i of at least $u_i(c_i^*(x) - \epsilon)$. No matter what the joint message and transfer of agents in state x , the total consumption, net of punishments, cannot exceed $c_1^*(x) + c_2^*(x)$. To achieve this net consumption both agents must report $r^1 = r^2 = 1$ and $\tilde{x}^1 = \tilde{x}^2 = x$. Any interim side contract satisfying (IC3) must induce zero punishment and, given that utility must be $u_i(c_i^*(x) - \epsilon)$, consumption must be at least $c_i^*(x) - \epsilon$. $\square$