

Price Equilibrium, Efficiency, and Decentralizability
in Insurance Markets with Moral Hazard

Richard Arnott
and
Joseph Stiglitz
(Stanford University)

Working Paper No. 254
January 1993

PRICE EQUILIBRIUM, EFFICIENCY, AND DECENTRALIZABILITY
IN INSURANCE MARKETS WITH MORAL HAZARD

Richard Arnott*

and

Joseph Stiglitz**

Revised version
January 1993

Preliminary draft: Please do not cite or quote without the permission of one of the authors.

*Department of Economics
Boston College
Chestnut Hill, MA 02167
USA

**Department of Economics
Stanford University
Stanford, CA 94305
USA

Abstract

In this paper, we investigate the descriptive and normative properties of competitive equilibrium with moral hazard when firms offer "price contracts" which allow clients to purchase as much insurance as they wish at the quoted price.

We show that a price equilibrium always exists and is one of three types:

- i) zero-profit price equilibrium — zero profit, zero effort, full insurance
- ii) positive-profit price equilibrium — positive profit, positive effort, partial insurance
- iii) zero-insurance price equilibrium — zero insurance, zero profit, positive effort.

Suppose a client purchases an additional unit of insurance from an insurer and consequently reduces accident-avoidance effort. This will lower the profitability of insurance the client has obtained from other insurers. In setting price, the insurer neglects this effect, however. Thus, price insurance entails an externality. We show under what circumstances this externality can be fully internalized by a linear tax on insurance sales.

Actual insurance contracts lie in the middle ground between exclusive quantity contracts (where an individual is effectively constrained to purchase all his insurance from one firm) and price contracts. We argue that our analysis of price contracts sheds light on the welfare properties of actual insurance contracts. Notably, since the externality we identify will still be operative, the taxation of insurance sales is typically desirable.

Price Equilibrium, Efficiency, and Decentralizability in Insurance Markets with Moral Hazard*

1. Introduction

During the past twenty years, there has been a growing awareness of moral hazard and related incentive problems. Surprisingly, however, there has been little analysis of the nature of equilibrium — the conditions under which equilibrium exists, and when it does what its properties are.

Broadly speaking, the literature has proceeded as follows. Pauly (1974) first analyzed price equilibrium with moral hazard; competitive insurance firms offer price contracts, allowing their clients to purchase as much insurance as they wish at the quoted price. He argued (not completely correctly, as we shall see) that since firms offer actuarially fair price contracts, individuals purchase full insurance and hence have no incentive to prevent the accident. It was then recognized that the moral hazard problem can be mitigated if individuals are rationed in the amount of insurance they can purchase at a particular price. One way this can be achieved is for firms to offer exclusive quantity contracts which specify both a price and quantity of insurance, and to prohibit their clients from obtaining additional insurance from other insurers, i.e. to insist that they be their clients' exclusive agent. Variants of this exclusive quantity equilibrium have been analyzed by Helpman and Laffont (1975), Shavell (1979), Stiglitz (1983), Prescott and Townsend (1984), and Arnott and Stiglitz (1987). As well, the exclusive quantity equilibrium has been employed extensively in the principal-agent literature, and its welfare properties have been investigated (Prescott and Townsend (1984), Arnott and Stiglitz (1986, 1989)).

An important unresolved issue is: What contract form will arise when exclusivity cannot be costlessly enforced? The literature has examined a number of different approaches. Hellwig (1983a) investigated endogenous communication between firms; Arnott and Stiglitz (1987) examined equilibrium with alternative assumptions concerning the set of admissible contracts; Kahn and Mookherjee (1989) focused on a coalition-

* The authors are indebted to the National Science Foundation, the Olin Foundation, the Hoover Institution, and the Social Sciences and Humanities Research Council of Canada for financial support.

proof Nash equilibrium in quantity contracts; Bizer and De Marzo (1992) examined a sequential equilibrium in which firms can condition their contracts on a client's past purchases of quantity contracts, but not on her subsequent purchases; and Arnott and Stiglitz (1991) examined a situation in which firms can restrict their clients' insurance purchases but not the insurance they obtain from non-market sources.

In this paper, we take a different tack. Actual insurance market equilibria lie in the middle ground between the exclusive contract equilibrium and the price equilibrium. Since, however, this middle ground contains a huge variety of possible contract forms, insight may be gained from thorough analysis of the extreme cases. The positive and normative properties of the exclusive contract equilibrium have been extensively investigated. The price equilibrium has, however, been neglected. Thus, in this paper, we investigate the descriptive and normative properties of price equilibrium in a competitive insurance market with moral hazard. We do not attempt a fully general analysis; rather, we provide a thorough analysis of a particularly simple case in which there is a single good and a single fixed-damage accident.

In his analysis of the price equilibrium, Pauly (1974) argued as follows. Since insurance firms offer actuarially fair price contracts, individuals purchase full insurance — that quantity of insurance which equalizes the marginal utility of income in the events "accident" and "no-accident." When the accident does not directly affect the utility function, so that equalization of marginal utilities implies equalization of total utilities, individuals have no incentive to prevent the accident. Thus, price equilibrium is characterized by zero profits, zero "effort" (at accident avoidance) and full insurance. Pauly also suggested that the price equilibrium entails the "overprovision" of insurance.

Pauly's analysis, while insightful, was not conclusive. Is it not possible that with zero effort the price of insurance would be so high that individuals would choose to purchase no insurance?¹ Is it obvious that equilibrium entails firms offering actuarially fair insurance? And does "overprovision" mean simply that, though infeasible, it would be beneficial to restrict insurance purchases so as to stimulate effort, or does it imply that government intervention, in some form, would be beneficial? In his positive analysis, Pauly erred in neglecting corner solutions and second-order conditions. Helpman and Laffont (1975) demonstrated that moral hazard can give rise to nonconvexities and that

¹Pauly recognized this possibility but did not analyze it.

these nonconvexities can cause non-existence of the equilibrium Pauly described. But they did not provide a full analysis of the price equilibrium.

We show that a price equilibrium always exists and is unique. It may have the characteristics that Pauly describes; we term such an equilibrium a zero-profit price equilibrium. But there are two other types of price equilibrium as well. In the first, individuals purchase no insurance; we term this a zero-insurance price equilibrium. In the second, insurance is sold but not at an actuarially fair rate; there is positive profit, partial insurance, and positive effort; this type of equilibrium we term a positive-profit price equilibrium. We also clarify in what sense a price equilibrium entails the overprovision of insurance. We assume that the government, like insurance firms, cannot observe how much insurance a specific individual purchases, but can monitor and hence tax firms' aggregate sales of insurance. With price insurance, there is an uninternalized externality. When it sells a client an extra unit of insurance which will cause her to reduce effort, a firm neglects that doing so will lower the profitability of the insurance she has obtained from other carriers. Thus, the private cost of providing insurance falls short of the social cost. At first glance, it might appear that this problem can be fully remedied by a linear "Pigouvian" tax on insurance.² The situation is, however, complicated by the nonconvexities to which moral hazard may give rise. We show that, depending on circumstances, the tax may be fully effective (as effective as direct rationing), partially effective, or ineffective. Government intervention is potentially justified since the market cannot mimic such a tax. We shall also argue that the linear taxation of insurance sales is potentially desirable with other contract forms whenever firms cannot perfectly monitor their clients' total insurance purchases.

Section 2 presents the basic model and the diagrammatic apparatus we shall employ. Price equilibrium in the absence of government intervention is analyzed in section 3. Section 4 investigates the linear taxation of price insurance. Section 5 provides some discussion, while section 6 concludes.

2. The Basic Model

²Helpman and Laffont (1975) recognized this inefficiency and also claimed (not completely correctly, as we shall see) that it can be corrected by taxation when indifference curves are convex.

We employ a very simple model in which individuals are identical (to abstract from adverse selection problems). There is a single, fixed-damage accident, the probability of which, p , is a function of the level of effort devoted to accident avoidance, e . Moral hazard arises because the insurer is unable to observe a client's effort, and is therefore unable to write insurance contracts contingent on it. Hence, insurance is based solely on whether or not the accident occurred. In the absence of insurance and in the event of no accident, an individual's consumption is w , while if there is an accident it is $w - d$, where d is the accident damage. An individual's insurance pays α (the "(net of premium) payout" or "benefit") if an accident occurs, and if it does not the insured pays β (the "premium"). Thus, the insurance the individual obtains is characterized by (α, β) . Also, where y_0 is consumption in the event of no accident, and y_1 consumption in the event of accident,

$$y_0 = w - \beta \quad y_1 = w - d + \alpha. \quad (1)$$

The individual's expected utility is

$$EU = (1 - p(e))U_0(y_0, e) + p(e)U_1(y_1, e). \quad (2)$$

We make the simplifying assumption that the expected utility function has the form³

$$EU = (1 - p(e))u(y_0) + p(e)u(y_1) - e. \quad (2')$$

This function is separable in consumption and effort, and is event-independent by which we mean that the accident does not affect the utility-from-consumption function — the accident does not cause pain, nor does it alter tastes. There is positive but diminishing marginal utility from consumption. Also, we measure effort so that $e = 0$ corresponds to minimum effort and assume that $p' < 0$ and $p'' > 0$ for $e > 0$, and $\lim_{e \downarrow 0} p(e) = p(0) \equiv \bar{p} < 1$.

We shall develop our analysis in $\alpha - \beta$ space (see Figure 1). We define $q \equiv \frac{\beta}{\alpha}$ to be the price of insurance⁴ corresponding to (α, β) (geometrically, it is the slope of a ray from the origin to (α, β)) and call α the quantity of insurance. Thus, a price insurance

³Some of the complications which arise when the event-contingent utility functions are not separable in consumption and effort are discussed in Arnott and Stiglitz (1988a).

⁴Later, we shall have occasion to work with price contracts which emanate from a point other than the origin. In such situations, the price of insurance is the slope of the price line.

contract at price q' , under which a client may purchase as much positive insurance as she wishes at price q' , is drawn as a ray from the origin with slope q' .

The individual with insurance (α, β) chooses effort to maximize utility, i.e.

$$\max_e EU = (1 - p(e))u_0 + p(e)u_1 - e \equiv V(\alpha, \beta) \quad (3)$$

where $u_0 \equiv u(w - \beta)$ and $u_1 \equiv u(w - d + \alpha)$. The first-order condition is

$$-(u_0 - u_1)p'(e) - 1 = 0 \quad \text{if } e > 0; \quad (4)$$

the marginal cost of effort is unity, and the marginal benefit of effort equals the reduction in the probability of accident from an extra unit of effort, $-p'$, times the increase in utility from a unit reduction in the probability of accident, $u_0 - u_1$. Under the assumption of separable utility, e is a continuous function of α and β , $e(\alpha, \beta)$. Furthermore, $\frac{\partial e}{\partial \alpha} < 0$ and $\frac{\partial e}{\partial \beta} < 0$ for $e > 0$; as more insurance is provided, the individual reduces effort. It is assumed that $e(0, 0) > 0$. Substituting $e(\alpha, \beta)$ into the expression for expected utility gives $V(\alpha, \beta)$, the indifference curves corresponding to which can be plotted in $\alpha - \beta$ space.

The individual's marginal rate of substitution between the premium and net payout is

$$\left. \frac{d\beta}{d\alpha} \right|_{\bar{V}} = -\frac{V_\alpha}{V_\beta} = \frac{u_1' p}{u_0' (1 - p)}. \quad (5a)$$

As more insurance is provided, diminishing marginal utility causes $\frac{u_1'}{u_0'}$ to decrease, but individuals take less care, and as a result $\frac{p}{1-p}$ increases. For this reason, indifference curves between net payout and premium will not in general be "convex"⁵ (see Figure 1):

$$\left. \frac{d^2\beta}{d\alpha^2} \right|_{\bar{V}} = \begin{cases} -s \left[(A_1 + sA_0) + \frac{u_1' (p')^3}{(1-p)^2 p p''} \right] & \text{for } e > 0 \\ -s(A_1 + sA_0) & \text{for } e = 0 \end{cases} \quad (5b)$$

⁵We say that indifference curves with the normal shape are "convex," even though in $\alpha - \beta$ space normally-shaped indifference curves are concave according to the definition.

where $s \equiv \frac{u_1' p}{u_0' (1-p)}$ and $A_i \equiv -\frac{u_i''}{u_i'}$, $i = 0, 1$.

Even though effort is a continuous function of (α, β) , an individual's purchases of insurance may not be a continuous function of the price of insurance because nonconvexity of the indifference curves may cause the price-consumption line (PC-line) to be discontinuous.⁶ This greatly complicates the analysis to follow.

The set of insurance contracts for which expected returns are non-negative,

$$\pi(\alpha, \beta) \equiv \beta(1 - p(e(\alpha, \beta))) - \alpha p(e(\alpha, \beta)) \geq 0, \quad (6)$$

is referred to as the feasibility set and denoted by F . The boundary of this set is referred to as the resource constraint or zero-profit locus (ZPL) (see Figure 1). With movement up the zero-profit locus, effort falls, the probability of accident increases, and to maintain zero profits the ratio $\frac{\beta}{\alpha}$ must increase. The slope of the zero-profit locus is

$$\left. \frac{d\beta}{d\alpha} \right|_{ZPL} = \begin{cases} \left[p - \frac{(\alpha + \beta)u_1'(p')^3}{p''} \right] / \left[(1-p) + \frac{(\alpha + \beta)u_0'(p')^3}{p''} \right] & \text{with } e > 0 \\ p(0)/(1-p(0)) & \text{with } e = 0. \end{cases} \quad (7)$$

It is shown in Arnott and Stiglitz (1988a) that the feasibility set is never convex. It is the lack of convexity of the feasibility set, combined with the lack of quasi-concavity of $V(\alpha, \beta)$, which give rise to many of the problems that are the concern of this paper.

We list some geometric properties that will be employed in the subsequent analysis.

1. The zero-effort line (ZEL) is the locus of (α, β) such that $(u_0 - u_1) \lim_{e \downarrow 0} p'(e) + 1 = 0$ (recall eq. (4)). Effort is zero everywhere above the ZEL. See Figure 2.
2. Because our assumption of separable utility implies that with $e > 0$ effort falls as α or β increases, below the ZEL the probability of accident rises monotonically along any ray from the origin.

⁶The PC-line is defined in the usual way as the locus of points of maximum utility on price lines (rays through the origin).

3. The full-insurance line (FIL) is $u_0' = u_1'$, which with event independence coincides with $\alpha + \beta = d$ and $u_0 = u_1$. From (4) if $\lim_{e \downarrow 0} p'(e) = -\infty$, then the FIL and the ZEL coincide; if, however, $\lim_{e \downarrow 0} p'(e)$ is finite, then the FIL lies outside the ZEL, the case depicted in Figure 2. The slope of indifference curves at the FIL is $p(0)/(1 - p(0)) \equiv q^*$.
4. At any point on the ZPL, such as K in Figure 1, the slope of the ray from the origin is $q_K \equiv \frac{\beta_K}{\alpha_K} = \frac{p_K}{1 - p_K}$.
5. The PC-line can intersect the ZPL at only one point with $\alpha, \beta > 0$. ($q = \frac{pu_1'}{(1-p)u_0'}$ and $q = \frac{p}{1-p}$ together imply $u_1' = u_0'$ and $q = \frac{p(0)}{1-p(0)}$, which intersect only once).
6. At the origin, the indifference curve through the origin is steeper than the ZPL (compare (5a) and (7)).
7. Since the numerator in (7) is always positive, there is only one point on the ZPL corresponding to a given β .

3. Price Equilibrium

In a price contract, the firm is required to sell to each client as much positive insurance as she wishes to buy at the price⁷ quoted in the contract. An individual may purchase from one or many firms; it makes no difference. There are many firms selling insurance and selling insurance entails no cost.

Individuals will purchase insurance at the lowest available price. Hence, equilibrium, if it exists, will be characterized by a single price, $\hat{q}$. Thus, $\beta = \hat{q}\alpha$, where q is chosen by firms and α by the individual. Faced with this price, individuals choose $\hat{\alpha} \geq 0$ such that

$$\hat{\alpha} = \arg \max_{\alpha} V(\alpha, \hat{q}\alpha). \quad (8)$$

⁷We do not consider random price contracts, where the random price is realized either before (ex ante) or after (ex post) the individual's effort decision. Arnott and Stiglitz (1988b) provides a discussion of randomization with moral hazard when firms offer exclusive quantity contracts.

The solution to (8) for all $q > 0$ characterizes the price–consumption line. For q such that $\alpha = 0$ (the price of insurance is so high that the individual would prefer to purchase no insurance), the origin is the point on the PC–line.

Firms must make non–negative profits in equilibrium. Thus, $\hat{q}$ must satisfy $(\hat{\alpha}\hat{q})(1 - \hat{p}) - \hat{\alpha}\hat{p} \geq 0$ which implies either $\hat{q} \geq \frac{\hat{p}}{1 - \hat{p}}$ for $\hat{\alpha} > 0$ or $\hat{\alpha} = 0$, where $\hat{p} \equiv p(e(\hat{\alpha}, \hat{q}, \hat{\alpha}))$.

It is easy to show that equilibrium exists, is unique, and occurs at the lowest point (corresponding to the lowest price) on the PC–line in $\alpha - \beta$ space consistent with non–negative profits.

3.1 The zero–profit price equilibrium

If the equilibrium is at an interior point (i.e. $\alpha > 0$, $\beta > 0$) and if zero profits are made, then from property 5 of the geometry of the problem, the equilibrium must be at the unique intersection of the FIL and ZPL, which we label E. With the form of the utility function assumed, full insurance implies zero effort. Thus, at such an equilibrium, there is zero profit, zero effort, and full insurance.

Pauly (1974) analyzed only the zero–profit price equilibrium, and it has generally been assumed that price equilibrium is always of this type. In fact, however, there are three mutually exclusive and collectively exhaustive types of price equilibrium, of which the zero–profit price equilibrium is one.

3.2 The zero–insurance price equilibrium

In this case, the PC–line lies everywhere outside the feasibility set except at the origin. At high prices no insurance is bought. As price is lowered, a price is reached where individuals switch discontinuously to purchasing so much insurance and expending so little effort that losses are made. Thus, the origin is the equilibrium. This case is illustrated in Figure 3. The zero–insurance price equilibrium is characterized by zero insurance, zero profit, and positive effort, and entails inactivity. Since this form of price equilibrium requires that there be a discontinuity in the PC–lines across the zero–profit locus, it can arise only when indifference curves are nonconvex.

3.3 The positive–profit price equilibrium

An equilibrium with positive profit is possible only if any insurance policy with a lower price loses money; thus, the price–consumption line must have a discontinuity across the zero–profit locus at the equilibrium point. This is illustrated in Figure 4. H is the (lower jump) point on the PC–line corresponding to the lowest price at which non–negative profits are made, q^H . Firms offering insurance at price q^H make positive profits. When, however, the price of insurance is lowered just slightly below q^H , individuals purchase much more insurance and take much less care, with the result that the insurance contract becomes unprofitable. Such an equilibrium is characterized by positive profit, positive effort, and partial insurance.⁸

This case occurs if some part of the PC–line lies inside the feasibility set and the PC–line does not intersect the ZPL, as illustrated in Figure 4. It also occurs if the PC–line does intersect the ZPL, but some portion of the PC–line with $q < q^*$ lies in F , as shown in Figure 5. This situation is peculiar since it requires that as price is lowered, at some price there is a discontinuous decrease in the quantity of insurance purchased. Arnott and Stiglitz (1988a) provides the condition under which this can occur.

Drawing the above results together, we have:

Proposition 1: When only price contracts are admitted, equilibrium always exists and is unique. It occurs at the lowest point on the price–consumption line consistent with non–negative profits (which may be the origin) and is one of the following three types.

1. Zero–profit price equilibrium. Zero profit, zero effort, full insurance. This occurs when and only when the PC–line intersects the zero–profit locus and no point on the price–consumption line with $q < q^*$ lies in the feasibility set.
2. Zero–insurance price equilibrium. Zero profit, positive effort, zero insurance. This occurs when and only when the price–consumption line

⁸Since effort is non–increasing in the amount of insurance offered, if effort is zero at H , it is zero beyond H . This implies that the indifference curve passing through H is convex beyond H , which is inconsistent with a jump discontinuity from H to H' (see Figure 4). Thus, effort is positive at H . Also, since H is on

the price–consumption line, $\left(\frac{d\beta}{d\alpha}\bigg|_V\right)^H = \left(\frac{u_1' p}{u_0'(1-p)}\right)^H = q^H > \frac{p^H}{1-p^H}$ (because of positive profits)

$\Rightarrow (u_1' - u_0')^H > 0 \Rightarrow$ partial insurance.

lies everywhere outside the feasibility set except at the origin, and as a result the market is inactive.

3. Positive-profit price equilibrium. Positive profit, positive effort, partial insurance. This occurs when and only when the conditions for the previous two equilibrium types are not satisfied.

3.4 Primitive conditions for the type of price equilibrium

If indifference curves are convex (see (5b) for the primitive condition in terms of the utility and probability-of-accident functions), the equilibrium is a zero-profit price equilibrium. The proof is straightforward. With convex indifference curves, the PC-line is continuous. This rules out both the zero insurance and positive-profit price equilibrium. If, to the contrary, indifference curves are concave everywhere in the feasibility set inside the full-insurance line, there is a zero-insurance price equilibrium since the origin is preferred⁹ to E and no interior point in the feasibility set is on the price-consumption line.

Recall that greater risk aversion makes indifference curves more convex, while stronger moral hazard (a higher elasticity of the probability of accident with respect to the amount of insurance provided) makes them more concave. Thus, if the probability of accident increases rapidly with the amount of insurance provided and if individuals are only slightly risk-averse (so that indifference curves are concave everywhere in the feasibility set when effort is positive) equilibrium will entail zero price insurance. If, contrarily, moral hazard is not severe and individuals are strongly risk-averse (so that indifference curves are convex) price equilibrium will entail full insurance.

When indifference curves are neither concave nor convex everywhere in the feasibility set where effort is positive, which equilibrium type occurs depends on global properties of the utility and probability-of-accident functions. In these circumstances, the primitive conditions for occurrence of the three equilibrium types are more difficult to characterize. We have, however, obtained some results.

⁹The slope of an indifference curve on the FIL is $\frac{p(0)}{1-p(0)}$. If indifference curves are concave below the FIL then the average slope of the indifference curve through the origin, from the origin to the FIL is less than $\frac{p(0)}{1-p(0)}$. This implies that its point of intersection with the FIL is below E , which in turn implies that utility is higher at the origin than at E .

We can bring together results that have previously been derived to provide partial characterizations of feasible equilibrium types. The first characterization is in terms of the point of maximal utility on the price line corresponding to zero effort, $q^* \equiv \frac{p(0)}{1-p(0)}$ (i.e., the point of intersection of the PC-line and the q^* -price line).¹⁰ Recall that a necessary condition for a zero-profit equilibrium is that E , which lies on the q^* -price line, be on the PC-line, and that a necessary condition for a zero-insurance equilibrium is that the point of maximal utility on the q^* -price line be at the origin. Thus, if the point of maximal utility is at the origin, a zero-profit equilibrium can be ruled out; if it is at E , a zero-insurance price equilibrium can be ruled out; and if it is in between, equilibrium must be of the positive-profit type (the situation depicted in Figure 6).

There are also local conditions which rule out particular equilibrium types. A local condition that rules out a zero-profit price equilibrium is illustrated in Figure 7. A necessary condition for a zero-profit equilibrium is that E be a local utility maximum on the q^* -price line. This is not possible if, as drawn in Figure 7, the indifference curve through E is concave just below E .¹¹ Also, if indifference curves are convex near the origin, the origin is not on the PC-line, and equilibrium cannot be of the zero-insurance type.

4. On the Desirability of the Linear Taxation of Price Insurance

Can the price equilibrium be improved upon by government intervention? This depends, of course, on what informational and contractual advantages the government has vis-à-vis the private sector. We shall assume that the government, like insurance firms, can observe neither an individual's effort nor her total insurance purchases. But it is able to monitor the aggregate sales of insurance by each company. One might call this an

¹⁰Points beyond E can be ruled out since the slope of an indifference curve along the price line q^* beyond E is $\frac{u_1' p(0)}{u_0' (1-p(0))} = \frac{u_1'}{u_0'} q^* < q^*$ (beyond the FIL).

¹¹Two necessary and sufficient conditions for the indifference curves through E to be concave just below E are that: i) $\lim_{e \downarrow 0} p' = -\infty$ (so that the ZEL and FIL coincide — geometric property 3) and ii)

$\lim_{e \downarrow 0} \frac{(p')^3}{p''} = -\infty$ (so that indifference curves are concave just below the ZEL — from (5b) — and hence just below the FIL, and hence just below E). This pair of conditions is satisfied if, for example $p(e) = \hat{p} - ke^\varepsilon$ for $\varepsilon < \frac{1}{2}$ and e sufficiently small.

informational advantage, but the same results obtain when this information is public. The government can then employ its monopoly on coercive taxation to tax insurance firms' sales. It could employ non-linear taxation, but in equilibrium firms would then sell that amount of insurance which minimized their average tax rate.¹² Thus, without loss of generality, we assume to simplify the analysis that the government applies a linear tax on firms' insurance sales.¹³

We first determine under what conditions the linear tax on firms' insurance sales is fully effective, i.e. achieves the optimum conditional on the observability of effort. Then, when this optimum cannot be so decentralized, we investigate the best that can be done.

4.1 The optimum conditional on the unobservability of effort

Return to Figure 1. The point θ is the point of maximum utility on the zero-profit locus. Imagine that only the government provides insurance. It is easy to check that Figure 1 is unaltered, except that the zero-profit locus becomes the government budget constraint or the economy's resource constraint. Thus, θ is the point of maximum utility conditional on the unobservability of effort.¹⁴ We shall refer to θ simply as "the optimum."

Why is the optimum not generally¹⁵ achievable under price insurance? Plot the zero-profit locus in $\alpha - q$ space. Adopting the convention that costs, like the premium, are payable only when no accident occurs, the ZPL may be viewed as the average cost curve for providing insurance coverage, $AC = \frac{p}{1-p}$. The corresponding marginal cost curve then relates the slope of the ZPL to α , while the corresponding demand curve relates the slope of the indifference curve to α , along the ZPL. Figure 8 depicts these loci for the case where the FIL and ZEL coincide, indifference curves are well-behaved, the demand curve is downward-sloping, and the ZPL is everywhere positively-sloped

¹²In the absence of taxation, since there are no scale effects, firm scale is indeterminate. With non-linear taxes, the scale of firms would be such as to minimize the average tax rate.

¹³The linearity of the tax on firms' insurance sales should not be confused with the linearity of the consumer's "insurance-possibility locus." Even if the government were to impose a non-linear tax on firms' insurance sales, individuals would still be able to purchase as much insurance as they wish at the equilibrium price — they would still face a linear insurance-possibility locus.

¹⁴A utility improvement over θ may be possible through randomization of insurance premia and payouts. See Arnott and Stiglitz (1988b). We ignore this possibility.

¹⁵The optimum can be at the point E , when effort is very onerous. When it is, zero taxation of insurance sales is optimal. We shall ignore this case for the remainder of the paper since we judge it to be practically unimportant.

and has a slope discontinuity at E . Then up to α_E , the average cost curve is upward-sloping and the marginal cost curve lies above the average cost curve. At E , there is a downward discontinuity in MC , and a slope discontinuity in AC . Beyond E , AC and MC coincide at a value of $\frac{p(0)}{1-p(0)}$. The demand curve intersects MC at θ , and AC at E . The social optimum occurs at θ , with insurance sales α_θ . If the corresponding demand price, $\hat{q}$, is charged, insurance firms operate at a profit, and will bid the price down. Equilibrium occurs at the price q_E since demand is satisfied and firms make zero profits. The optimum can be decentralized by imposing the "Pigouvian" tax $\tau^* = \hat{q} - q_\theta$.

An analogy might be helpful. Suppose that an increase in the number of tomatoes (units of insurance coverage) a person buys increases the cost of producing each of the tomatoes he buys by the same amount. Whether the person buys tomatoes from one or many sellers, competitive equilibrium will entail $p = AC$. The person will purchase the optimal number of tomatoes when $p = MC$, but at this price profits are positive. The optimum can be achieved by applying a tax on tomatoes equal to $MC - AC$, evaluated at the optimum.

4.2 Decentralization of the optimum

With the linear taxation of insurance firms' aggregate sales, the consumer price of insurance is defined naturally as the ratio of the premium payable to the insurance company in the no-accident event to the net payout from the insurance company in the event of accident. The tax revenue is redistributed to individuals as a lump-sum subsidy (though we admit the possibility that part of it may be destroyed if doing so improves expected utility). Thus,

$$y_0 = w - q\alpha + \ell \quad y_1 = w - d + \alpha + \ell, \quad (9)$$

where ℓ is the lump-sum transfer. With taxation, the natural space to work in is $(\alpha + \ell, q\alpha - \ell)$ (see Figure 9). Note that zero insurance purchased corresponds to the point $(\hat{\ell}, -\hat{\ell})$ in this space.

The major result of this sub-section is

Proposition 2: A pair of sufficient conditions under which the optimum can be attained when firms provide price insurance and there is linear taxation of aggregate insurance sales is:

- i) indifference curves be convex; and
- ii) $(1 - \hat{p})\hat{q} - \hat{p} > (1 - \tilde{p})\tilde{q} - \tilde{p}$ on the $\hat{\ell}$ -price-consumption line $\forall \tilde{q} < \hat{q}$, where $\hat{p}$ is the probability of accident at the optimum, $\hat{q}$ the consumer price of insurance at the optimum, $\hat{\ell}$ the lump-sum subsidy at the optimum, and the $\hat{\ell}$ -price-consumption line the price-consumption line with $(\hat{\ell}, -\hat{\ell})$ as origin.

Condition ii) states that insurance firms' pre-tax profits must be greater at price $\hat{q}$ than at any lower price. We shall first assume convexity of indifference curves and demonstrate that with convexity, condition ii) is both necessary and sufficient for decentralizability of the optimum with linear taxation. We shall then show why nonconvexity of indifference curves may prevent decentralization of the optimum via taxation of insurance sales.

4.2.1 Decentralization with convex indifference curves

To start, we set up an artificial planning problem and then investigate decentralization. We imagine that the planner is able to choose the price of insurance, but not the quantity, and redistributes any profits through a lump-sum subsidy to consumers, ℓ . In this case, the government's problem is to

$$\begin{aligned}
 \max_{q, e, \alpha, \ell} \quad & EU = u(w - q\alpha + \ell)(1 - p(e)) + u(w - d + \alpha + \ell)p(e) - e \\
 \text{s.t. i)} \quad & q\alpha(1 - p(e)) - \alpha p(e) \geq \ell \\
 \text{ii)} \quad & (e, \alpha) = \arg \max_{e, \alpha} (EU; q, \ell),
 \end{aligned} \tag{10}$$

where i) is the resource constraint, and ii) the constraints imposed by the individual's choosing how much effort to expend and how much insurance to purchase, taking both the price of insurance and the lump-sum payment as given. We have put an inequality constraint on i) since we permit the government to destroy resources.

Turn to Figure 9. Draw in the line tangent to the optimum point θ and extend it back until it intersects the 45°-line in the lower quadrant. The point of intersection, O' , is $(\hat{\ell}, -\hat{\ell})$, where $\hat{\ell}$ is the optimal lump-sum transfer, and the slope of the line, $\hat{q}$, gives the optimal price of insurance. If the individual is given a lump-sum transfer $\hat{\ell}$ and allowed

to purchase as much insurance as he wants at the price $\hat{q}$, he will choose the optimum point θ . Furthermore, since θ is on the resource constraint, the government breaks even on the provision of insurance, i.e. the revenue from insurance premia just covers insurance payouts plus lump-sum transfers.

This result by itself is not of great practical interest since it is hard to envision a situation where the government is the sole provider of insurance and is restricted to price insurance. But it does suggest a decentralization mechanism for which the optimum may be achievable. The mechanism is as follows: The government taxes total insurance purchases at rate $\hat{t}$ and distributes a lump-sum subsidy $\hat{\ell}$ to each consumer. $\hat{t}$ and $\hat{\ell}$ are set such that: a) firms break even if the economy is at the optimum; and b) the government's budget balances if the economy is at the optimum. Now, a firm's profits per unit of insurance are

$$\pi = (1 - p)q - p - t. \quad (11)$$

Thus, condition a) is that

$$(1 - \hat{p})\hat{q} - \hat{p} - \hat{t} = 0. \quad (12a)$$

And condition b) is that

$$\hat{\ell} = \hat{t}\hat{\alpha}, \quad (12b)$$

where $\hat{\alpha}$ is the level of α at the optimum (per (10)).

We say that θ is decentralizable with the tax system $(\hat{\ell}, \hat{t})$ iff: a) taking $(\hat{\ell}, \hat{t})$ as given, competition among firms results in an equilibrium price of insurance of $\hat{q}$; and b) given $\hat{\ell}$ and $\hat{q}$, individuals choose to buy a quantity of insurance $\hat{\alpha}$. The "if" part of the definition follows from the construction, and the "only if" part is straightforward.

Condition b) follows from the convexity of the indifference curves. It might appear that condition a) is necessarily satisfied since at the optimum firms make zero profits. The equilibrium price cannot be higher than $\hat{q}$. Suppose it were. Then a firm which offered insurance at price $\hat{q}$ would capture the entire market and not make a loss—contradiction. But is it necessarily the case that firms will make negative profits at all lower prices of insurance? If a firm lowers the consumer price from $\hat{q}$ to say $\tilde{q}$, the individual will choose the point of maximum utility on the price line having slope $\tilde{q}$ and

origin $(\hat{\ell}, -\hat{\ell})$. We term the locus of points so generated the $\hat{\ell}$ -price-consumption line. Thus, condition ii) of Proposition 2 is the requirement that the firm make negative profits at all lower prices than $\hat{q}$.

Now,

$$\left(\frac{d\pi}{dq} \right)_{\hat{\ell}-PC} = \left((1-p) - (q+1)p' \frac{de}{dq} \right)_{\hat{\ell}-PC}, \quad (13)$$

where $\hat{\ell}-PC$ denotes "evaluated along the $\hat{\ell}$ -price-consumption line." Ordinarily, one expects that a decrease in the price of insurance will discourage effort and increase the probability of accident, in which case $\left(\frac{d\pi}{dq} \right)_{\hat{\ell}-PC} > 0$. But a decrease in the price of insurance can, with decreasing absolute risk aversion, have such a strong income effect that individuals purchase less insurance; and this effect can be strong enough that the increase in the individual's effort induced by the purchase of less insurance causes the probability of accident to fall by so much that profits increase (the conditions under which this occurs are derived in Appendix 1). The reduction in the price of insurance below $\hat{q}$ puts the economy outside the resource constraint. Since the government is paying out $\hat{\ell}$ per individual and receiving $\hat{t}\alpha$ per individual in taxes, with $\hat{\ell}$ and $\hat{t}$ fixed, the individual's reduced purchases causes the government to lose money. This effect can be so strong that the government's loss exceeds the economy's loss — firms make a profit (see Appendix 1).

Helpman and Laffont (1975) overlooked this possibility. They consequently concluded incorrectly that convexity of indifference curves is a sufficient condition for decentralization of the optimum in this context.

The possibility that insurance firms' profits increase when the price of insurance is lowered strikes us as decidedly abnormal; hence, we describe the situation where condition ii) of Proposition 2 is satisfied as the normal case.

We have established that in the normal case and with convex indifference curves, the optimum can be decentralized through the linear taxation of insurance sales. The resulting equilibrium entails a two-part tariff, a lump-sum subsidy for participation in the market plus a linear tax on insurance. This cannot be achieved by the market alone, since if a single firm were to offer a contract $(\hat{\ell}, \hat{q})$ an individual would purchase an infinitesimal amount of insurance from the firm so as to receive the lump-sum subsidy,

and purchase the bulk of his insurance from other firms, which, since they do not provide a subsidy, can offer insurance at a lower price. Thus, through its coercive powers of taxation, the government essentially induces the market to offer a two-part tariff which the market cannot do on its own.

4.2.2 Nonconvexity of indifference curves

Nonconvexity of indifference curves provides another reason why the optimum may not be decentralizable using the linear taxation of insurance. This is shown in Figure 10; at the price of insurance corresponding to the optimum, the individual would choose to purchase insurance characterized by G rather than θ . The optimum may, however, be decentralizable even when the indifference curve through θ is nonconvex. What matters is whether θ is on the $\hat{\ell}$ -price-consumption line. Thus, Proposition 2 may be strengthened and simplified:

Proposition 2': Necessary and sufficient conditions under which the optimum can be attained when firms provide price insurance and there is linear taxation of aggregate insurance sales are:

- i) The optimum lie on the $\hat{\ell}$ -price-consumption line.
- ii) $(1 - \hat{p})\hat{q} - \hat{p} > (1 - \tilde{p})\tilde{q} - \tilde{p}$ on the $\hat{\ell}$ -price-consumption line $\forall \tilde{q} < \hat{q}$.

4.2.3 Interpretation

In the introduction, we indicated that the linear tax on insurance sales can be regarded as a Pigouvian tax. Here we expand on this interpretation. Consider a marginal insurance firm that sells an individual an increment of insurance, which supplements the insurance he has purchased from other firms. Since it provides only an infinitesimal amount of insurance, it takes the probability of accident as given. It ignores that its sale will cause the individual to reduce her effort, which will adversely affect the profitability of the individual's other insurance. This is the externality. Now turn back to Figure 8. In terms of that diagram, the marginal insurer sets his price equal to the average rather than the marginal cost of insurance. The Pigouvian prescription is to impose a tax, τ^* , equal to the difference between the marginal and average cost, evaluated at the optimum. Since, per Figure 8, the Pigouvian tax is paid only when an accident does not occur, whereas the tax on insurance sales is paid whatever the outcome, we wish to show that $(1 - \hat{p})\tau^* = \hat{t}$:

$$\begin{aligned}
(1 - \hat{p})\tau^* &= (1 - \hat{p})(MC - AC)_\theta \\
&= (1 - \hat{p})\left(\frac{d\beta}{d\alpha}\bigg|_{ZPL} - \frac{p}{1 - p}\right)_\theta \\
&= (1 - \hat{p})\left(\hat{q} - \frac{\hat{p}}{1 - \hat{p}}\right) = \hat{t} \quad (\text{from (12a)}).
\end{aligned}$$

Thus, Proposition 2' indicates the conditions under which a Pigouvian tax on insurance permits decentralization of the optimum.

This externality interpretation is helpful. But since the same equilibrium occurs when an individual purchases all his insurance from a single firm, the interpretation is incomplete. An alternative intuition, mentioned earlier, is that linear pricing and zero profits force the firm to average-cost price, whereas efficiency requires marginal-cost pricing.

4.3 Ameliorative taxation when the optimum cannot be decentralized

The above discussion raises two related questions. If the optimum is not decentralizable through linear taxation when individual insurance purchases are unmonitorable, what is the best that can be done — what is the constrained optimum? And is taxation at some rate always better than no taxation at all?

4.3.1 Determination of the constrained optimum

We treat the former question first. To answer it, we construct a decentralizability locus. A point is on the decentralizability locus if, for a given lump-sum subsidy $\tilde{\ell}$, it is the lowest point on the $\tilde{\ell}$ -price-consumption line, $(\tilde{\alpha} + \tilde{\ell}, \tilde{q}\tilde{\alpha} - \tilde{\ell})$, for which there exists a tax rate on insurance $\tilde{t}$ such that:

- i) the government at least balances its budget, $\tilde{t}\tilde{\alpha} \geq \tilde{\ell}$;
- ii) insurance firms at least break even, $(1 - \tilde{p})\tilde{q} - \tilde{p} - \tilde{t} \geq 0$; and
- iii) insurance firms cannot increase profits by lowering the price of insurance.

This is the point of maximum utility associated with $\tilde{\ell}$ such that individuals do not wish to change their purchases of insurance, firms have no incentive to alter price, and both

insurance firms and the government at least break even. The constrained optimum is then the maximum utility point on the decentralizability locus. In Figure 11, the point M is on the decentralizability locus as long as conditions i) – iii) are satisfied. Unfortunately, condition iii) does not have a neat geometric characterization. Note that the constrained optimum may entail insurance firms making positive profits and/or the government destroying some tax revenue; in the former case, the constrained optimum is defined attaching zero welfare weight to such profits.¹⁶

We summarize the above results in:

Proposition 3: The constrained optimum is the point of maximum utility on the decentralizability locus. If the optimum point, θ , is on the decentralizability locus, then the optimum can be decentralized through the linear taxation of firms' insurance sales when firms offer price insurance; otherwise, it cannot be.

The information required to solve for the constrained optimum is substantial. Thus, it is of interest to investigate under what circumstances a welfare-improving tax can be found with less information. We first investigate when a small tax is desirable.

4.3.2 When a small tax is welfare-improving

Since a zero-insurance price equilibrium is a corner solution, a small tax will generally not induce individuals to purchase insurance and therefore has no effect on welfare. In contrast, a small tax applied to a positive-profit price equilibrium is always welfare improving. This is illustrated in Figure 12. The positive-profit price equilibrium in the absence of taxation is at the point H . Now imagine raising the consumer price of insurance to q_I , holding utility fixed. The consumer's insurance purchases shift to I . The decentralization of I entails a positive lump-sum transfer, ℓ_I , and hence a tax on insurance (since $t\alpha_I \geq \ell_I$). But I cannot be an equilibrium, since there are lower points on the ℓ_I -price-consumption line consistent with non-negative government revenue and non-negative profits for insurance firms. The equilibrium for ℓ_I must be at the lower jump point on the ℓ_I -PC line. At this point firms do not have an incentive to lower price. If they did so, not only would the economy operate at a loss, but also the government would operate at a surplus since insurance sales would increase.

¹⁶The analysis would be similar if partial or full welfare weight were attached to firms' profits.

The analysis of the effect of a small tax applied at a zero-profit price equilibrium is considerably more complex. To simplify the problem somewhat, we assume that the profit per unit of insurance falls when the price of insurance is lowered (The analysis in Appendix 1 implies that a sufficient condition for this is that $\frac{-u''}{(u')^2}$ be increasing in y).

Under our assumption of separable, event-dependent utility, full insurance is provided at the zero-profit equilibrium. As well, effort is zero. A small tax on firms' insurance sales has two opposing effects on welfare. The tax results in the individual receiving less than full insurance, which has a negative second-order welfare effect. The tax may also stimulate effort, which has a positive welfare effect. Which effect dominates depends on the properties of the probability-of-accident function as effort approaches zero, since this determines the order of magnitude of the effort-stimulating effect on welfare.

It turns out that the desirability of a small tax on insurance at a zero-profit price equilibrium depends on whether the zero-effort line coincides with the full-insurance line

— recall from geometric property 3 in section 2 that the two coincide if $\lim_{e \downarrow 0} -p'$ is infinite, and do not coincide otherwise. It also matters whether there is a slope discontinuity in the zero-profit locus across the zero-effort line; there is if $\lim_{e \downarrow 0} \frac{(p')^3}{p''}$ is non-zero, and there is not if the limit is zero. We noted earlier that if $\lim_{e \downarrow 0} p' = -\infty$ and $\lim_{e \downarrow 0} \frac{(p')}{p''} = -\infty$, then indifference curves just below the candidate zero-profit equilibrium point are concave, implying that the price equilibrium cannot be of the zero-profit type. This leaves us with five relevant cases, all of which are possible,¹⁷ depending on the properties of the probability-of-accident function:

¹⁷Case I: $p(e) = \bar{p} - ke^{\frac{1}{2}}$; case II: $p(e) = \bar{p} - ke^{\varepsilon}$, $1 > \varepsilon > \frac{1}{2}$; case III: $p(e) = \bar{p} - k_0e + k_1e^{\beta}$, $\beta > 2$; case IV: $p(e) = \bar{p} \exp(-ze)$; case V: $p(e) = \bar{p} - k_0e + k_1e^{\beta}$, $\beta \in (1, 2)$.

		$-\lim_{e \downarrow 0} \frac{(p')^3}{p''}$		
		infinite	finite	0
$-\lim_{e \downarrow 0} p'$	infinite	n.a.	I	II
	finite	III	IV	V

Let us start with case I, which is drawn in Figure 13. The algebra for this case is presented in Appendix 2. From geometric property 3, when $\lim_{e \downarrow 0} p' = -\infty$, the ZEL and the FIL coincide. From (7), it follows that there is a slope discontinuity in the zero-profit locus at E , with a discontinuous increase in the slope from below to above E . And from (5a), the indifference curves have a continuous slope as they cross the zero-effort line. Since the indifference curve through E , V_0 , is tangent to the ZPL "just above" E , it follows from the geometry of the problem that utility is higher at a point slightly below E on the ZPL, such as E' , than at E . It remains to show that E' can be decentralized with a small positive tax. The algebra is presented in Appendix 2. Here we shall present a heuristic argument. Since the PC-line through E intersects V_1 above the ZEL, and since the indifference curves are locally convex, it is necessary to apply a positive tax on insurance to attain the point E' . If this tax, $t_{E'}$, is applied, the government balances its budget by providing a lump-sum subsidy $\ell_{E'} = \alpha_{E'} t_{E'}$, and since the economy is on the zero-profit locus, insurance firms break even. Thus, a small tax unambiguously increases welfare. It has a negative second-order effect on welfare through exposing the individual to more risk via less-than-full insurance, but a positive first-order effect on welfare through its stimulation of effort.

In case III, shown in Figure 14, and in cases IV and V, the full-insurance line lies strictly outside the zero-effort line. In these cases, a small tax on insurance generates a negative second-order welfare effect by disequalizing the event-contingent marginal utilities and a zero effect on effort, and is therefore unambiguously harmful. The reader can picture case II by considering Figure 13, but without the slope discontinuity in the zero-profit locus at E . A small tax on insurance is desirable if the indifference curve V_0

has less curvature just below E than the ZPL,¹⁸ and is undesirable otherwise. This is the case where both the risk and effort–stimulating welfare effects of a small tax are second–order.

We summarize the above results in:

Proposition 4: The welfare effects of a small tax on firms' insurance sales (with the revenue redistributed to consumers in lump–sum fashion) depends on the nature of the price equilibrium. Such a tax is never welfare–improving at a zero–insurance price equilibrium, but always welfare–improving at a positive–profit price equilibrium. At a zero–profit price equilibrium, the welfare effects of the tax depend on the curvature properties of the probability–of–accident function as effort approaches zero.

4.3.3 When there is a large welfare–improving tax

A large welfare–improving tax exists if any part of the decentralizability locus lies on a lower indifference curve (corresponding to higher utility) than the price equilibrium. In this subsection, we briefly investigate the circumstances under which this will be the case.

If the price equilibrium is a positive–profit price equilibrium, there is almost always a large welfare–improving tax. The argument is essentially the same as that for the desirability of a small tax for this type of equilibrium.

In the case of a zero–profit price equilibrium, there may not be a large, welfare–improving tax. An example where there is not is shown in Figure 15. V_0 is the utility at the zero–profit price equilibrium E , and V_1 the utility at the optimum. The hatched area is the set of feasible, utility–improving insurance allocations, relative to E . Within this area, the indifference curves between V_0 and V_1 are concave. As a result, no feasible, utility–improving insurance allocations can be decentralized with a price line. For the same reason, there may also be no large, welfare–improving tax with a zero–insurance price equilibrium.

We conclude this subsection with

¹⁸From (7), the curvature of the ZPL in the positive effort region contains terms in p''' . Since there are no natural restrictions on the sign or magnitude of p''' , either the ZPL or V_0 can have greater curvature.

Proposition 5: Whether there is a large, welfare-improving tax depends on the nature of the nonconvexities, and hence on global rather than local properties of the utility and probability-of-accident functions. A sufficient condition for the nonexistence of a welfare-improving tax is that indifference curves be concave in the region of feasible, utility-improving (relative to the price equilibrium) insurance allocations. One sufficient condition for the existence of a welfare-improving tax is that the price equilibrium be of the positive-profit type.

5. Discussion

5.1 Potential advantages of government intervention

Whenever one encounters a model in which there is scope for welfare-improving government intervention, one should ask: What is it that the government can do that private agents cannot? Relatedly, has the model unduly restricted the actions of private agents or given the government unreasonably broad powers? And, finally, has the model been consistent in its treatment of information or has it artificially provided the government with informational advantages over the market?

We first consider whether the government is informationally-advantaged vis-à-vis the market, and then consider whether it is contractually-advantaged.

5.1.1 Informational advantages

The government, unlike the collectivity of firms, can require that all insurance policies be registered. The advantage that the government has over the collectivity of firms in this context derives from its monopoly on legislation and the legal enforcement of contracts. It does not need to employ compulsion, since it can declare illegal or refuse to enforce unregistered policies.

The government may utilize the information acquired from registration in a variety of ways: it can nationalize the provision of insurance; it can regulate the market, by for example imposing a ceiling on each individual's total insurance purchases; it can impose non-linear taxation of individuals' insurance purchases, essentially forcing individuals to choose insurance packages from along the zero-profit locus; or it can make information on each individual's insurance purchases available to insurance companies.

The registration of sales is particularly effective in the context of insurance. Since insurance is intrinsically non-transferable, the problem of secondary markets does not arise, as it would for example if the government were to require the registration of cigarette sales in an attempt to ration each individual's cigarette consumption.

But there are presumably forms of insurance for which the registration of individual policies is excessively costly, or is deemed to constitute an invasion of privacy. Then the government has a more limited informational advantage over the market. It can employ its powers of legal compulsion to require that insurance firms report their aggregate sales of insurance. This is the informational situation considered in this paper. By itself this informational advantage does not justify government intervention since if the government were only to collect such information and make it public, the market equilibrium would be unaltered. It is this informational advantage along with the power to tax aggregate insurance sales that creates the potential for Pareto-improving government intervention.

5.1.2 Contractual advantages

The government has one obvious contractual advantage over the market. Through the nationalization of insurance or the regulation of a private monopoly insurer, it can effect contractual forms that may not be sustainable in an unregulated market; in particular, either institution results in exclusive contracts which may not be enforceable at reasonable cost in competitive equilibrium without government intervention. Here the advantages of the government over the market stems from its powers of proscription.

There are well-known problems with both nationalized industries and regulated monopolies. Thus, it is of interest to enquire whether the government can employ its unique powers to support contractual forms within the market that would be unsustainable without government intervention. As we have seen, this is precisely what the linear taxation of insurance achieves. The market, with the help of the government but not without it, is able to offer insurance in the form of a two-part tariff, a subsidy $\hat{\ell}$ for participating in the market, plus a unit price of insurance $\hat{q}$. This is not as efficient as having each individual face a non-linear price schedule, but is the best that can be done given that each individual's total insurance purchases are unobservable. Here the

contractual advantages of the government over the market stem from its monopoly on the coercive power of taxation.¹⁹

5.2 On the desirability of taxing insurance

In this paper we have investigated the conditions under which the taxation of price insurance is potentially desirable when moral hazard is present. When indifference curves are convex, taxation is generally desirable; indeed, the optimum conditional on the unobservability of effort can normally be achieved. When indifference curves are nonconvex, the case for taxation needs to be qualified. But the conditions under which there is no scope for Pareto-improving taxation are sufficiently odd that it seems safe to say that, even with nonconvexities, taxation at some rate is likely desirable.

There are, however, potential problems with the taxation of price insurance. The first is an empirical one; the informational requirements for determining the optimal tax rate would be substantial. It would not generally be known a priori whether indifference curves are convex. With nonconvex indifference curves the benefit from taxation is not generally a continuous function of the tax rate. And even with convex indifference curves, the benefit from taxation is not generally a concave function of the tax rate.²⁰ Second, in practical insurance situations moral hazard and adverse selection are both present. Thus, the effect of the linear taxation of insurance sales on the adverse selection problem would need to be taken into account.²¹

Despite these qualifications, we have uncovered a general principle: With price insurance and moral hazard, individuals can usually be made better off by the linear taxation of insurance sales. Because of an insurance externality, firms collectively over-insure their clients. This can be corrected, at least partially, by setting the price of insurance above the actuarially fair level, which can be achieved via taxation.

¹⁹Taxation has another role to play when moral hazard is present. In Arnott and Stiglitz (1986) we argued that shadow prices deviate from market prices when moral hazard is present. One may regard moral hazard as a distortion, associated with which is a deadweight loss (relative to the corresponding economy with symmetric information). The consumption of any good alters the magnitude of this deadweight loss; for example, the consumption of a fire extinguisher likely reduces the deadweight loss associated with fire insurance. The second best is achieved by setting the commodity tax on a good equal to the increase in deadweight loss from a unit increase in the consumption of the good.

²⁰If the FIL lies outside the ZEL a small tax on insurance is harmful even when the optimum can be decentralized with a tax at the right rate.

²¹Insurance loading would also need to be taken into account in practical situations. Even though the administrative costs associated with a policy are largely fixed, loading is often set proportional to the amount of insurance purchased, in which case it resembles a linear tax on insurance sales.

Earlier we discussed the middle ground between price insurance contracts and exclusive quantity contracts. We conjecture that in most if not all of this middle ground, a Pareto improvement can be made through the linear taxation of insurance sales.²² Consider, for example, one situation treated in Arnott and Stiglitz (1987) where insurers cannot observe their clients' total insurance purchases but can offer non-exclusive quantity contracts. This situation can be examined by reference to Figure 9. With no government intervention, the optimum cannot be sustained by the non-exclusive quantity contract $(\alpha_\theta, \beta_\theta)$ since other firms can profit by offering small supplementary contracts at a price between $\left(\frac{p}{1-p}\right)_\theta$ and $\left(\frac{u'_1 p}{u'_0(1-p)}\right)_\theta$. But if the government imposes a tax on insurance sales a rate $\hat{t}$ (per section 4.1), the optimum may be sustainable. The firm would offer the non-exclusive quantity contract $(\hat{\alpha}, \hat{\beta})$ where $\hat{\beta} = \hat{\alpha}\hat{q}$. This, combined with the lump-sum subsidy $\hat{\ell}$ from the government, would provide the individual with the optimal insurance package. If θ is on the decentralizability locus, there are no profitable, supplementary contracts. To break even with taxation, a supplementary contract would have to have a consumer price at least as high as $\hat{q}$. But since, by construction, individuals are purchasing their most-preferred quantity of insurance at price $\hat{q}$, there would be no demand for such supplementary contracts. There is also no profitable replacement contract. Thus, the optimum is sustainable with non-exclusive quantity contracts and the linear taxation of aggregate insurance sales if θ is on the decentralizability locus.

We therefore conjecture that the principle we have uncovered applies not only to price insurance but also to the middle-ground contracts that are actually offered in the market.

5.3 Principal-agent problems

The correspondence between principal-agent problems (with uncertainty and risk-averse agents) and insurance problems is well-recognized. Thus, everything we have said applies as well to principal-agent problems, though the specifics vary from one application to the next. In all the applications we can think of, principals try to enforce

²²Note, however, that the linear taxation of insurance sales will not discourage the provision of informal insurance since such insurance escapes the tax. In fact, with informal insurance, the taxation of insurance sales may be harmful by inducing a substitution away from formal towards less efficient informal insurance.

exclusivity but without complete success: employers discourage moonlighting; landlords require that their sharecroppers work no other land; banks deny credit to prospective clients with excessive indebtedness. Since exclusivity is never perfect, agents are likely to overinsure, in which case there is scope for Pareto-improving taxation of insurance "sales." In many situations, however, most notably the labor market, it is difficult to measure the insurance implicit in the contract, in which case taxation may be impractical.

But there is one important application in which the linear taxation of the implicit insurance is practical and quite likely desirable — credit markets. Bizer and De Marzo (1992) provide a persuasive model of credit markets in which individuals acquire loans sequentially and on each loan application list prior²³ debt, and in which prior debt is senior. That last incremental competitive loan provides insurance at the actuarially fair price; the lender ignores the effect of his loan on the riskiness of prior debt. With a fixed-damage "accident" this situation is similar to the price insurance analyzed in this paper.²⁴ Correspondingly, a linear tax on the insurance implicit in loans mitigates the moral hazard problem. There is a widespread perception that consumer credit, at least in the form of credit cards, is too easy to acquire. This perception may correspond to the over-insurance problem we have identified.

6. Concluding Comments

This paper provided an analysis of the positive and normative properties of price equilibrium in competitive insurance markets with moral hazard. In a price equilibrium, insurance firms offer price contracts which allow clients to purchase as much insurance as they wish at the quoted price. We demonstrated that a price equilibrium always exists and is of one of three types:

- i) zero-profit price equilibrium — zero profit, zero effort, full insurance.
- ii) positive-profit price equilibrium — positive profit, positive effort, partial insurance.

²³Banks may attempt to restrict future indebtedness during the loan period. Such exclusivity-like provisions are difficult to enforce, particularly in the context of L.D.C. debt (Kletzer (1984), and Eaton, Gersovitz, and Stiglitz (1986)).

²⁴With a fixed-damage accident, there are two classes of lenders. The more senior lenders who face no default risk provide a pure loan. The junior lenders lose their entire loan in the event of default. Their loans are part pure loan, part insurance. Seniority provisions seem uncommon in pure insurance markets, though they do occur in mortgage insurance.

- iii) zero–insurance price equilibrium — zero insurance, zero profit, positive effort.

And we identified the circumstances under which each of the three obtains. The possibility of the zero–insurance and positive–profit price equilibria stems from the nonconvexities (of indifference curves in the relevant space) to which moral hazard can give rise.

We showed that a price equilibrium is generally not Pareto efficient, conditional on the unobservability of effort, and enquired whether a utility improvement can be made by taxing insurance. We assumed that neither insurance firms nor the government can observe an individual's total purchases of insurance, but that firms' aggregate sales of insurance are public information. In these circumstances, the government can do no better than to impose a linear tax on insurance sales, redistributing the proceeds in lump–sum fashion.

We found that there is scope for utility–improving linear taxation of price insurance. In some cases, the optimum (conditional on the unobservability of effort) can be attained through the linear taxation of insurance; in others, a utility improvement can be achieved through taxation, even though the optimum cannot, and in yet other cases, taxation is ineffectual. We derived the conditions under which each of these results obtains. These conditions are complex and relate to global rather than local properties of the probability–of–accident and utility functions.

The price equilibrium is a special case in two respects. First, the informational assumptions underlying it are extreme; firms have no information on their clients' supplementary insurance. Second, the contractual assumptions are extreme; firms offer only price contracts. Nevertheless, we argued that many of the insights from this extreme case carry over to situations in which firms have imperfect information on their clients' supplementary insurance and in which the form of contract is unrestricted. Thus, we have established a general theoretical case for the linear taxation of insurance sales.

An intuition for the desirability of taxation is that insurance firms generate a negative externality through their sale of insurance. When a firm sells a client an extra unit of insurance, she lowers her effort which reduces the profitability of the insurance she has obtained from other insurers. Under our assumptions, the government has no informational advantage over the collectivity of firms, but through its monopoly on

taxation is able to exploit the available information in a way that the market by itself cannot.

Before a practical case can be made for the linear taxation of insurance sales, it will be necessary to extend the analysis to account for adverse selection, administrative loading, etc. and then to determine empirically whether the market failure we have identified is quantitatively important in specific contexts.

Bibliography

- Arnott, R.J., and J.E. Stiglitz, "Moral Hazard and Optimal Commodity Taxation," Journal of Public Economics 29, 1986, 1–24.
- _____, "Equilibrium in Competitive Insurance Markets with Moral Hazard," discussion paper 4, John M. Olin Program for the Study of Economic Organization and Public Policy. Princeton University, 1987.
- _____, "The Basic Analytics of Moral Hazard," Scandinavian Journal of Economics 90, 1988a, 383–413.
- _____, "Randomization with Asymmetric Information," Rand Journal of Economics 19, 1988b, 344–362.
- _____, "Moral Hazard and Nonmarket Institutions: Dysfunctional Crowding Out or Peer Monitoring?", American Economic Review 81, 1991, 179–90.
- Bizer, D. and P. De Marzo, "Sequential Banking," Journal of Political Economy 100, 1992, 41–61.
- Eaton, J., M. Gersovitz, and J.E. Stiglitz, "Pure Theory of Country Risk," European Economic Review 30, 1986, 481–513.
- Greenwald, B. and J.E. Stiglitz, "Externalities in Economies with Imperfect Information and Incomplete Markets," Quarterly Journal of Economics 101, 1986, 229–264.
- Hellwig, M.F., "Moral Hazard and Monopolistically Competitive Insurance Markets," 1983a, University of Bonn discussion paper #105.
- _____, "On Moral Hazard and Non-price Equilibrium in Competitive Insurance Markets," 1983b, University of Bonn discussion paper #109.
- Helpman, E., and J.-J. Laffont, "On Moral Hazard in General Equilibrium," Journal of Economic Theory 10, 1975, 8–23.
- Kahn, C., and D. Mookherjee, "Decentralized Exchange and Efficiency under Moral Hazard," 1989, mimeo.

- Kletzer, K.M., "Asymmetries of Information and LDC Borrowing with Sovereign Risk," Economic Journal 94, 1984, 287–307.
- Pauly, M., "Overprovision and Public Provision of Insurance: The Roles of Adverse Selection and Moral Hazard," Quarterly Journal of Economics 88, 1974, 44–62.
- Prescott, E., and R. Townsend, "Pareto Optima and Competitive Equilibria with Adverse Selection and Moral Hazard," Econometrica 52, 1984, 21–45.
- Rothschild, M., and J.E. Stiglitz, "Equilibrium in Competitive Insurance Markets: An Essay on the Economics of Imperfect Information," Quarterly Journal of Economics 90, 1976, 629–49.
- Shavell, S., "On Moral Hazard and Insurance," Quarterly Journal of Economics 93, 1979, 471–500.
- Stiglitz, J.E., "Risk, Incentives, and Insurance: The Pure Theory of Moral Hazard," The Geneva Papers on Risk and Insurance 8, 1983, 4–33.

Appendix 1

Analysis of $\left(\frac{d\pi}{dq}\right)_{\hat{\ell}-PC}$

The individual's choice problem is

$$\max_{e, \alpha} L = (1 - p(e)) u(w - \alpha q + \hat{\ell}) + p(e) u(w - d + \alpha + \hat{\ell}) - e. \quad (\text{A1.1})$$

The first-order conditions are

$$e: -p'(u_0 - u_1) - 1 = 0 \quad (\text{A1.2a})$$

$$\alpha: -(1 - p)qu_0' + pu_1' = 0. \quad (\text{A1.2b})$$

After some algebraic manipulations, one obtains

$$\left(\frac{de}{dq}\right)_{\hat{\ell}-PC} = \frac{1}{\Delta} \left(p' \alpha u_0' u_1' p \left(A_1 - \frac{u_1'}{u_0'} A_0 \right) \right), \quad (\text{A1.3})$$

where

$$\Delta = \begin{vmatrix} \frac{p''}{p'} & \frac{p' u_1'}{1 - p} \\ \frac{p' u_1'}{1 - p} & -(1 - p) q u_0' (A_0 q + A_1) \end{vmatrix} > 0 \quad (\text{A1.4})$$

from the second-order conditions for a maximum.

Note that $\left(\frac{de}{dq}\right)_{\hat{\ell}-PC} < 0$ requires that $A_1 - \frac{u_1'}{u_0'} A_0 > 0$. This holds if $\frac{-u''}{(u')^2}$ is

decreasing in y , which implies that absolute risk aversion be decreasing in y . Note also, since (using (5b))

$$\Delta = -\frac{p'' p u_1' C}{p q} \quad (\text{A1.5})$$

where $C \equiv \frac{d^2\beta}{d\alpha^2}\bigg|_{\bar{v}}$ is the degree of convexity of an indifference curve, that convexity of indifference curves is not only consistent with $\left(\frac{d\pi}{dq}\right)_{\hat{\ell}-PC} < 0$, but is locally a necessary condition for a maximum of the individual's choice problem.

Combining (13), (A1.3), and (A1.5) gives

$$\left(\frac{d\pi}{dq}\right)_{\hat{\ell}-PC} = (1-p) - \frac{m\alpha u_0' q(q+1)}{C} \left(A_1 - \frac{u_1'}{u_0'} A_0 \right), \quad (\text{A1.6})$$

where $m \equiv \frac{\partial e}{\partial(u_0 - u_1)} = -\frac{(p')^3}{p''}$ is a measure of the severity of moral hazard. It can be seen that $\left(\frac{d\pi}{dq}\right)_{\hat{\ell}-PC} < 0$ is "more likely" negative, the more severe is moral hazard, the less convex are indifference curves, and the more absolute risk aversion is decreasing.

Appendix 2

Analysis of $\lim_{t \downarrow 0} \frac{dV}{dt}$ at E when $\lim_{e \downarrow 0} p' = -\infty$ and $\lim_{e \downarrow 0} \frac{(p')^3}{p''}$ finite.

The individual's utility function is

$$V = (1 - p(e))u(w - \alpha q + \ell) + p(e)u(w - d + \alpha + \ell) - e. \quad (\text{A2.1})$$

The individual chooses α and e taking q and ℓ fixed. The first-order conditions are

$$e: -p'(e)(u_0 - u_1) - 1 = 0 \quad \text{or} \quad e = 0 \quad (\text{A2.2})$$

$$\alpha: -(1 - p)qu_0' + pu_1' = 0. \quad (\text{A2.3})$$

Budget balance for the government implies

$$t\alpha = \ell. \quad (\text{A2.4})$$

And zero profits for insurance firms implies

$$(1 - p)q - p - t = 0 \quad \text{or} \quad q = \frac{p + t}{1 - p}. \quad (\text{A2.5})$$

Substitution of (A2.4) and (A2.5) into (A2.2) and (A2.3) yields

$$e: -p'(e) \left(u \left(w - \frac{\alpha p(1+t)}{1-p} \right) - u(w - d + \alpha(1+t)) \right) - 1 = 0 \quad \text{or} \quad e = 0 \quad (\text{A2.2}')$$

$$\alpha: -(p + t)u' \left(w - \frac{\alpha p(1+t)}{1-p} \right) + pu'(w - d + \alpha(1+t)) = 0 \quad (\text{A2.3}')$$

Equations (A2.2') and (A2.3') characterize equilibrium as a function of t , incorporating budget balance for the government and zero profits for insurance firms. These two conditions imply that the equilibrium lies on the ZPL (now more appropriately termed the resource constraint or feasibility frontier). Thus, from total differentiation at (A2.2') and (A2.3'), one can determine the movement of equilibrium along the ZPL as t changes.

We shall proceed slightly differently. Define $b \equiv \frac{\alpha p(1+t)}{1-p}$ and $a = \alpha(1+t)$. The former is the total premium — the premium paid to insurance companies less the lump—

sum transfer from the government. The latter is the total net payout — the net payout from insurance companies plus the lump-sum transfer from the government. The way we shall proceed is to prove that the left-side derivative of $\left. \frac{dV}{db} \right|_{E, ZPL} < 0$, and then that for

$b < b_E$ the right-side derivative of $\left. \frac{db}{dt} \right|_{ZPL} < 0$. Together these results will establish the right-side derivative of $\left. \frac{dV}{dt} \right|_{E, ZPL} > 0$.

1. To prove that $\lim_{b \uparrow b_E} \left. \frac{dV}{db} \right|_{ZPL} < 0$.

The analysis of indifference curves and the ZPL from section 2 carry through, with a taking the place of α and b taking the place of β . Since we know from geometric property 7 that for any value of β , there is a unique point on the ZPL, then we may without ambiguity take the left-hand limit of $\left. \frac{dV}{db} \right|_{ZPL}$.

$$\begin{aligned} \lim_{b \uparrow b_E} \left. \frac{dV}{db} \right|_{ZPL} &= \left(\frac{\partial V}{\partial b} + \frac{\partial V}{\partial a} \left(\lim_{b \uparrow b_E} \left. \frac{da}{db} \right|_{ZPL} \right) \right) \Big|_E \\ &= \left[-u_0'(1-p) + u_1' p \left(\frac{(1-p)-k}{p+k} \right) \right] \Big|_E \quad (\text{from (7)}), \\ &= -\frac{ku'}{p+k}. \end{aligned} \quad (\text{A2.6})$$

where $k \equiv \left[-(a+b)u' \lim_{e \downarrow 0} \frac{(p')^3}{p''} \right] > 0$. Thus, movement down the ZPL just below E has a positive first-order effect on utility.

2. To prove that $\lim_{t \downarrow 0} \left. \frac{db}{dt} \right|_{ZPL} < 0$ for $b < b_E$.

Now, $b = \frac{\alpha p(1+t)}{1-p}$. Equations (A2.2') and (A2.3') give $e = e(t)$ and $\alpha = \alpha(t)$.

Thus,

$$\begin{aligned} \frac{db}{dt} &= \frac{\partial b}{\partial t} + \frac{\partial b}{\partial e} \frac{de}{dt} + \frac{\partial b}{\partial \alpha} \frac{d\alpha}{dt} \\ &= \frac{\alpha p}{1-p} + \frac{\alpha(1+t)}{(1-p)^2} p' \frac{de}{dt} + \frac{p(1+t)}{1-p} \frac{d\alpha}{dt}. \end{aligned} \quad (\text{A2.7})$$

Below E on the ZPL, both (A2.2') with $e > 0$ and (A2.3') apply. Thus, on this section of the ZPL

$$\begin{aligned}
 & \begin{bmatrix} \frac{p''}{p'} + (p')^2 u_0' \frac{\alpha(1+t)}{(1-p)^2} & p' \left(u_0' \frac{p(1+t)}{1-p} + u_1'(1+t) \right) \\ -p' u_0' + p' u_1' + (p+t) u_0'' \frac{\alpha(1+t)}{(1-p)^2} p' & (p+t) u_0'' \frac{p(1+t)}{1-p} + p u_1''(1+t) \end{bmatrix} \begin{bmatrix} de \\ d\alpha \end{bmatrix} \\
 &= \begin{bmatrix} -p \left(u_0' \frac{\alpha p}{1-p} + u_1' \alpha \right) \\ u_0' - (p+t) u_0'' \frac{\alpha p}{1-p} - p u_1'' \alpha \end{bmatrix} dt \tag{A2.8}
 \end{aligned}$$

In the limit as $t \downarrow 0$, $b \uparrow b_E$, $u_0' = u_1'$, and $u_0'' = u_1''$. Hence, in the limit as $t \downarrow 0$, these equations simplify to

$$\begin{bmatrix} \frac{p''}{p'} + \frac{(p')^2 u' \alpha}{(1-p)^2} & \frac{p' u'}{1-p} \\ \frac{p u'' p' \alpha}{(1-p)^2} & \frac{p u''}{1-p} \end{bmatrix} \begin{bmatrix} de \\ d\alpha \end{bmatrix} = \begin{bmatrix} -p \frac{u' \alpha}{1-p} \\ u' - \frac{p u'' \alpha}{1-p} \end{bmatrix} dt \tag{A2.8'}$$

Thus, in the limit

$$\begin{aligned}
 \Delta &= \frac{p'' p u''}{p'(1-p)} + \frac{(p')^2 u' \alpha p u''}{(1-p)^3} - \frac{(p')^2 u' \alpha p u''}{(1-p)^3} \\
 &= \frac{p'' p u''}{p'(1-p)} > 0 \tag{A2.9}
 \end{aligned}$$

And

$$\begin{aligned}
 \frac{d\alpha}{dt} &= \frac{1}{\Delta} \begin{vmatrix} \frac{p''}{p'} + \frac{(p')^2 u' \alpha}{(1-p)^2} & \frac{-p' u' \alpha}{1-p} \\ \frac{p u'' p' \alpha}{(1-p)^2} & u' - \frac{p u'' \alpha}{1-p} \end{vmatrix} \\
 &= \frac{1}{\Delta} \left(-\alpha \Delta + \left(\frac{p''}{p'} + \frac{(p')^2 u' \alpha}{(1-p)^2} \right) u' \right) \tag{A2.10}
 \end{aligned}$$

$$\begin{aligned}
\frac{de}{dt} &= \frac{1}{\Delta} \left| \begin{array}{cc} \frac{p'u'\alpha}{1-p} & \frac{p'u'}{1-p} \\ u' - \frac{pu''\alpha}{1-p} & \frac{pu''}{1-p} \end{array} \right| \\
&= \frac{1}{\Delta} \left(-\frac{(u')^2 p'}{1-p} \right)
\end{aligned} \tag{A2.11}$$

Substituting (A2.10) and (A2.11) into (A2.7) gives that in the limit as $t \downarrow 0$

$$\lim_{t \downarrow 0} \frac{db}{dt} \Big|_{ZPL} = \frac{p}{1-p} \left(\frac{1}{\Delta} \right) \frac{p''}{p'} u' = \frac{u'}{u''} < 0. \tag{A2.12}$$

Finally, combining (A2.6) and (A2.12) yields

$$\begin{aligned}
\lim_{t \downarrow 0} \frac{dV}{dt} \Big|_{ZPL} &= \left(\lim_{b \uparrow b_E} \frac{dV}{db} \Big|_{ZPL} \right) \left(\lim_{t \downarrow 0} \frac{db}{dt} \Big|_{ZPL} \right) \\
&= \left(\frac{-ku'}{p+k} \right) \left(\frac{u'}{u''} \right) = \left(\frac{k}{p+k} \right) \left(-\frac{(u')^2}{u''} \right) > 0.
\end{aligned} \tag{A2.13}$$

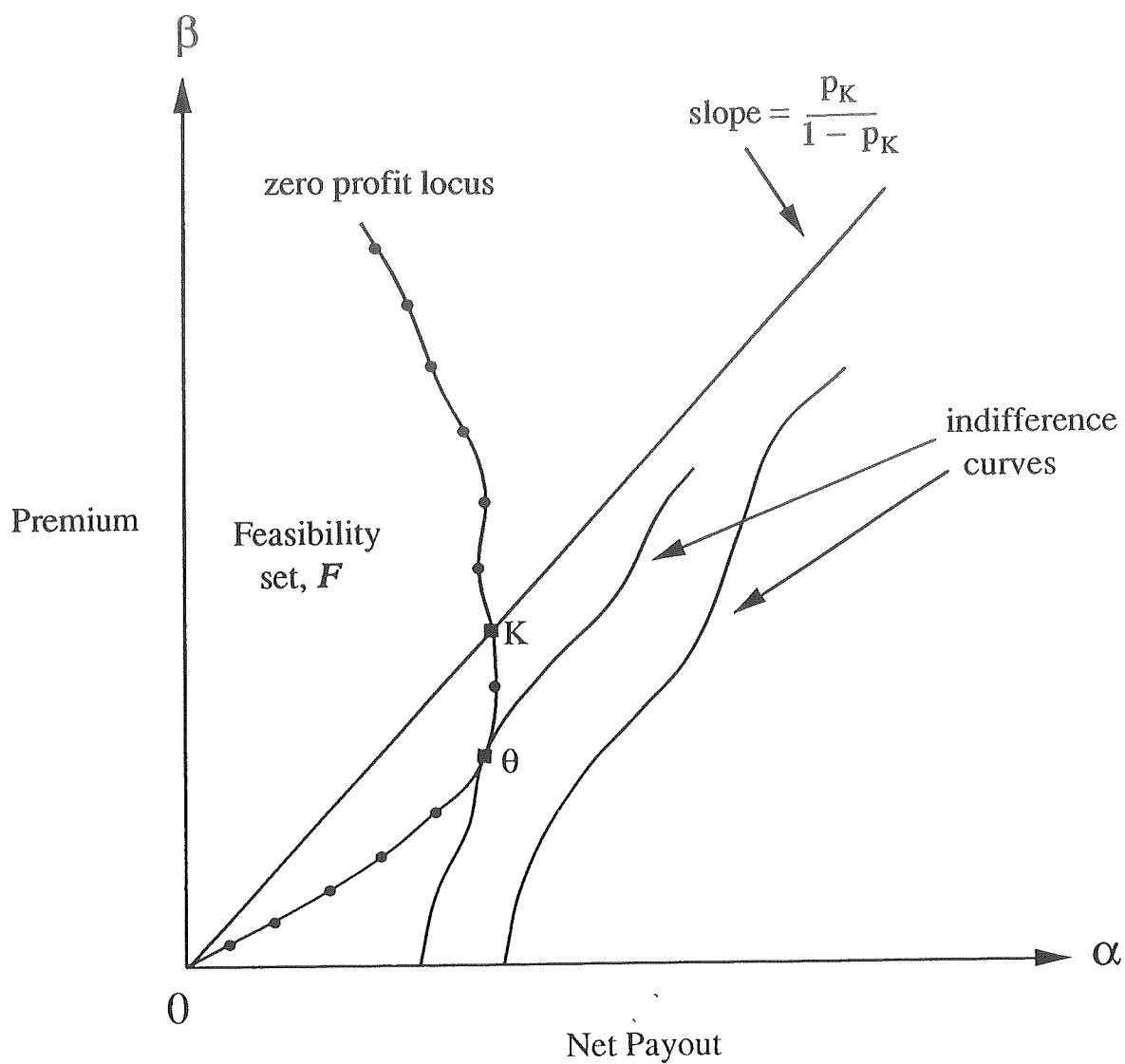


Figure 1: Basic diagram 1, continuum of effort levels

- (i) the zero profit locus is $\beta (1 - p) - \alpha p = 0$
- ii) the feasibility set is never convex
- iii) indifference curves are not necessarily convex.)

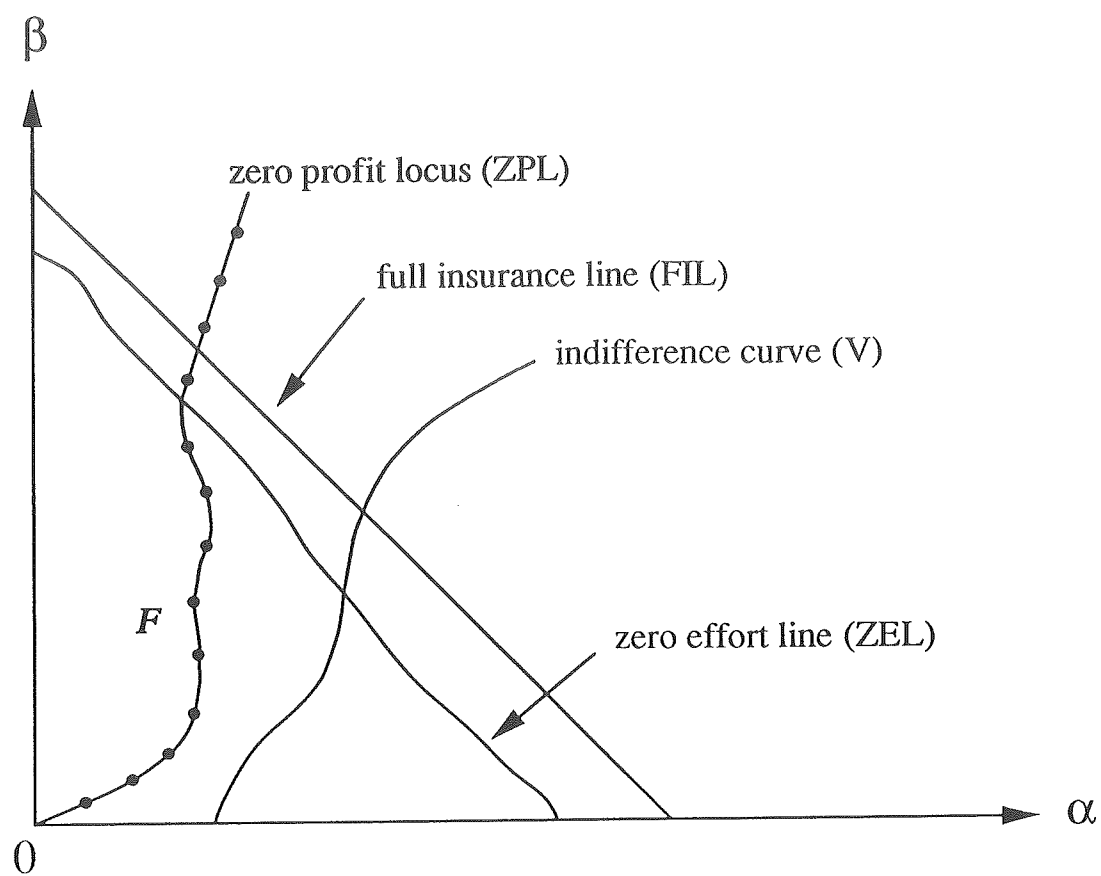


Figure 2: Basic diagram 2, $\lim_{e \downarrow 0} p'(e)$ finite

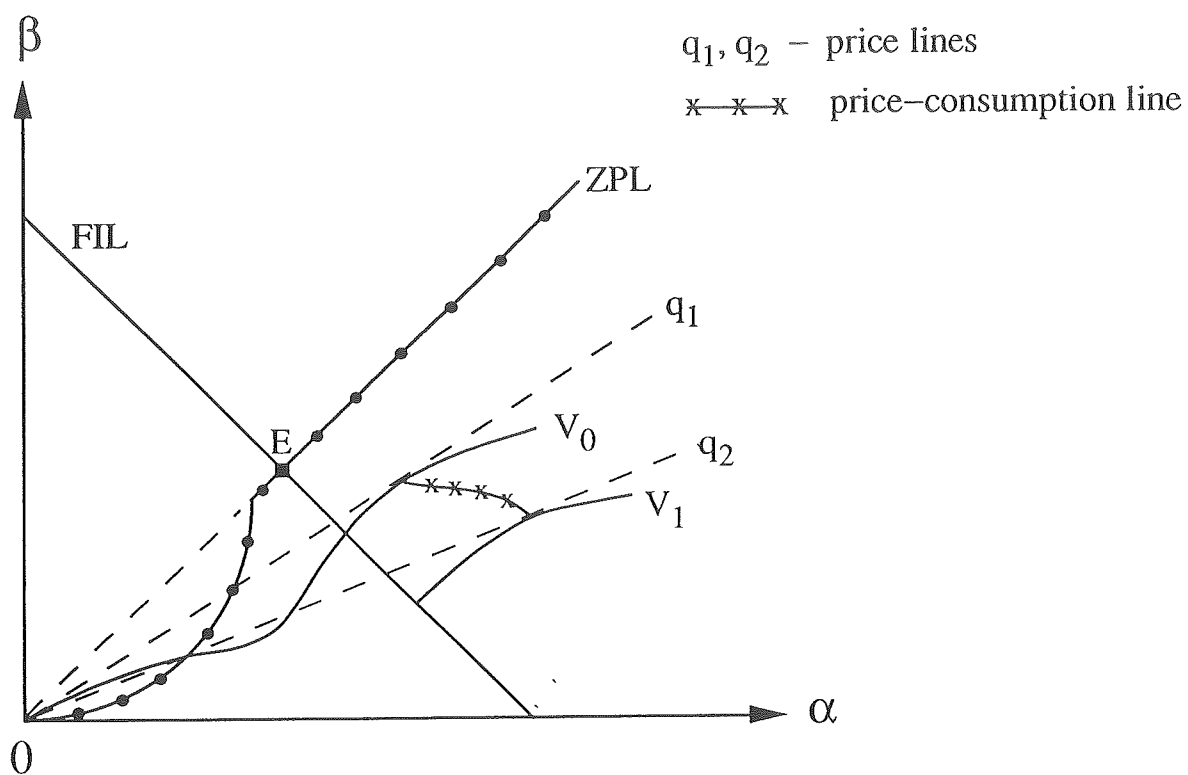


Figure 3: Zero insurance price equilibrium.
 (Observe that the individual's utility
 is higher at 0 than at E.)

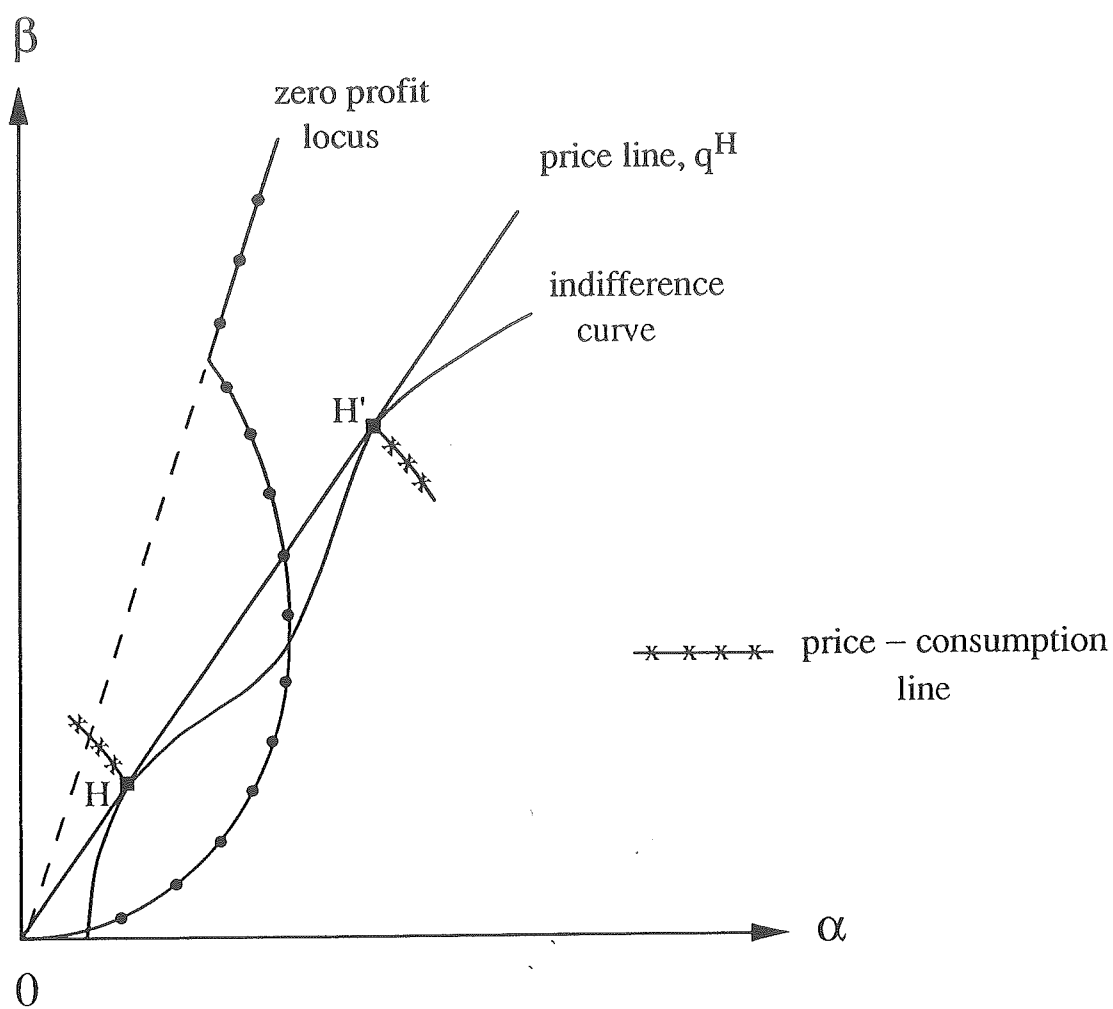


Figure 4: The price – consumption line can be discontinuous across the zero profit locus

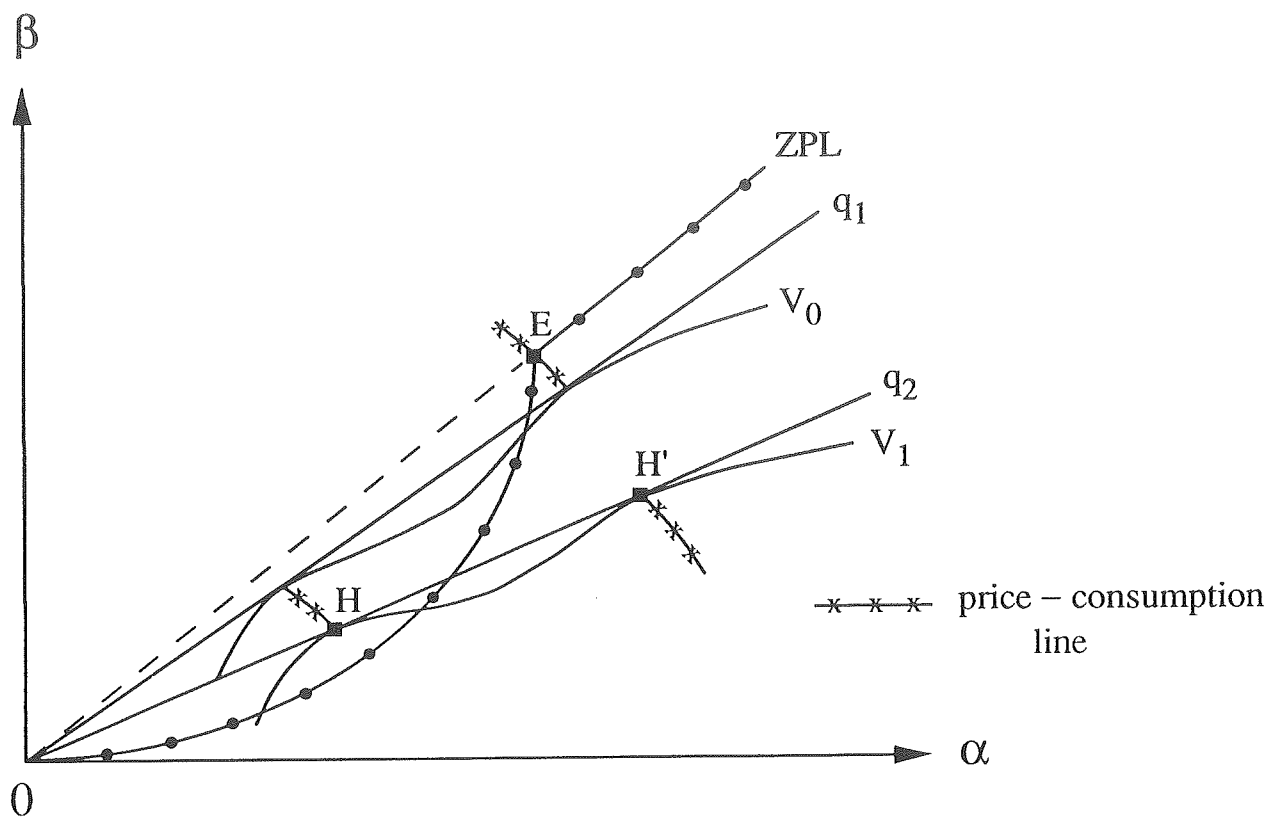


Figure 5: The price – consumption line intersects the ZPL but some portion of it with $q < q^*$ lies in F

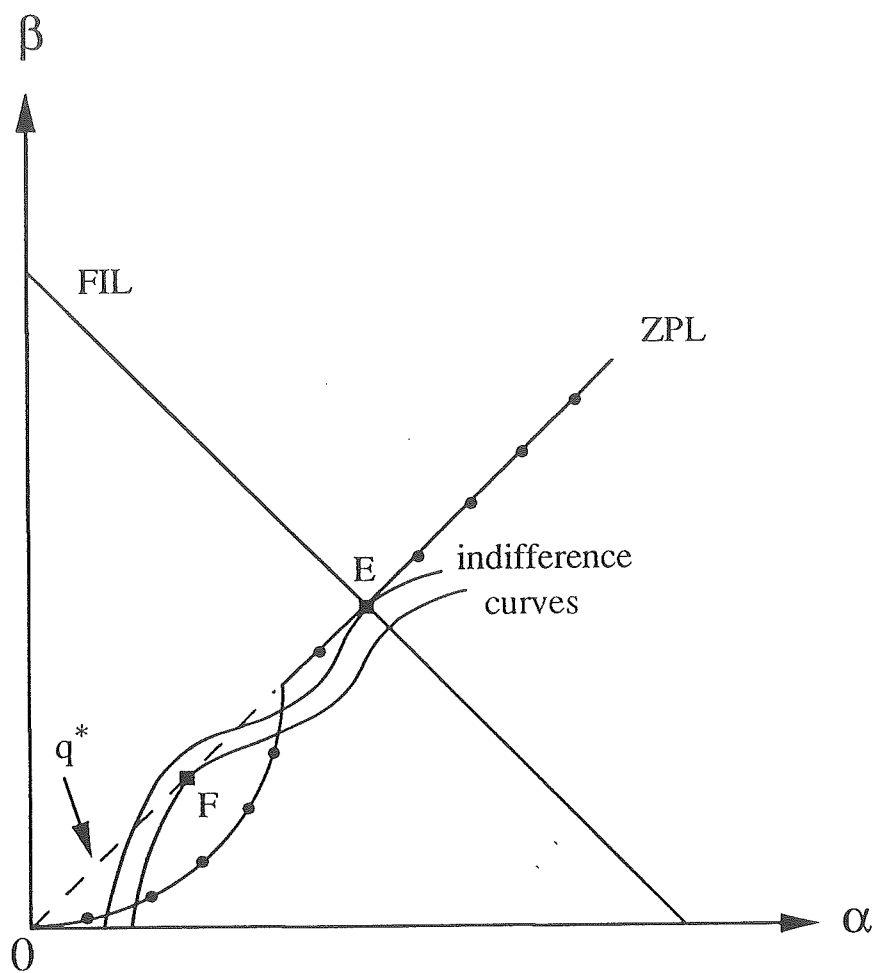


Figure 6: With this configuration, equilibrium is of the positive profit type

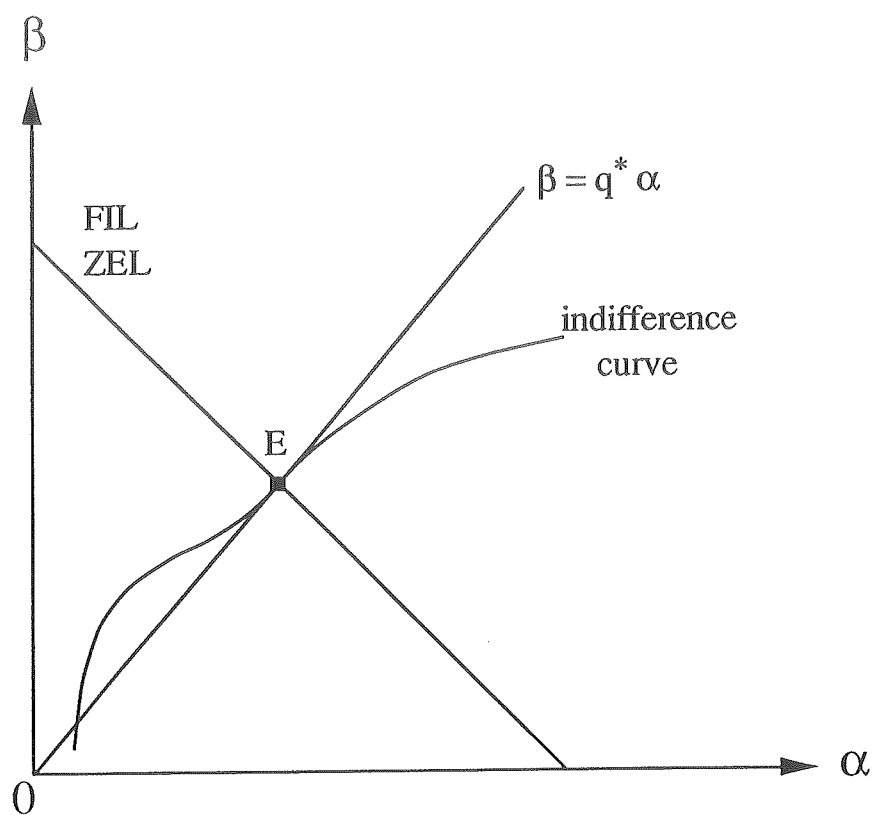


Figure 7: A local condition under which the price – consumption line does not intersect the zero profit locus

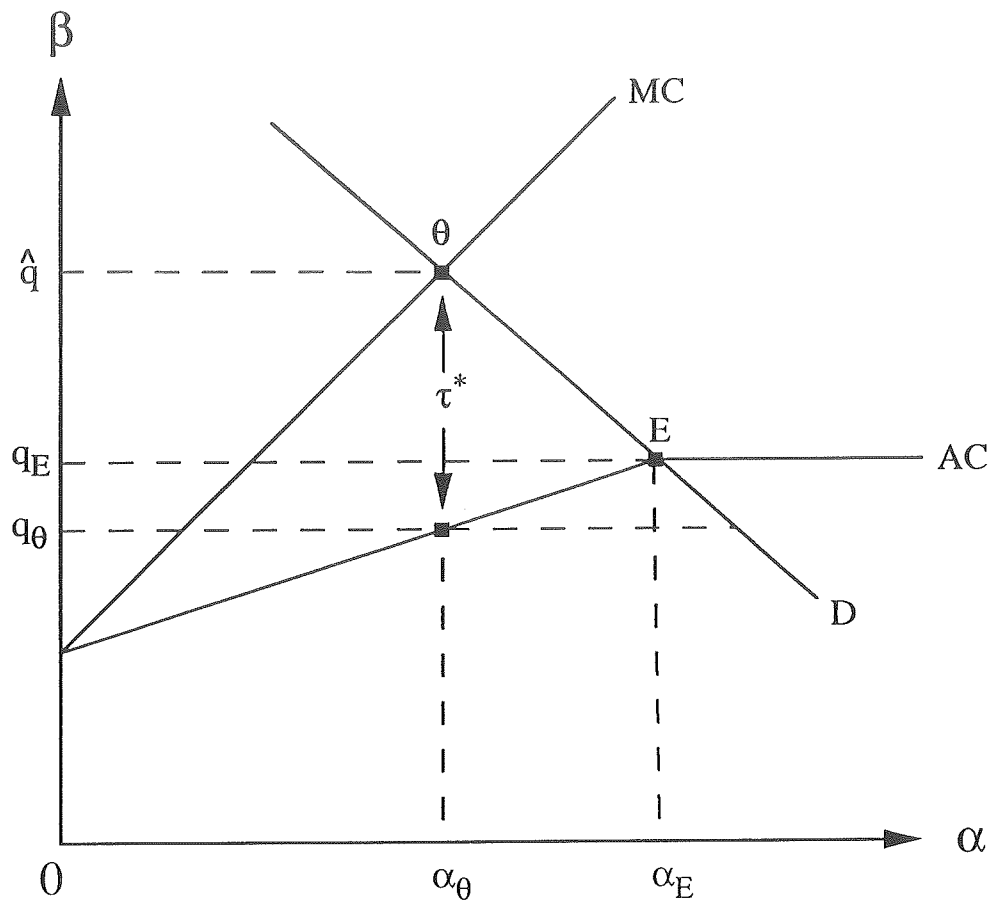


Figure 8: Why the optimum is not generally achievable under price insurance

(Note: drawn for the case where the FIL and ZEL coincide, indifference curves are well-behaved, the demand curve is downward-sloping, and the ZPL is everywhere positively-sloped and has a slope discontinuity at E)

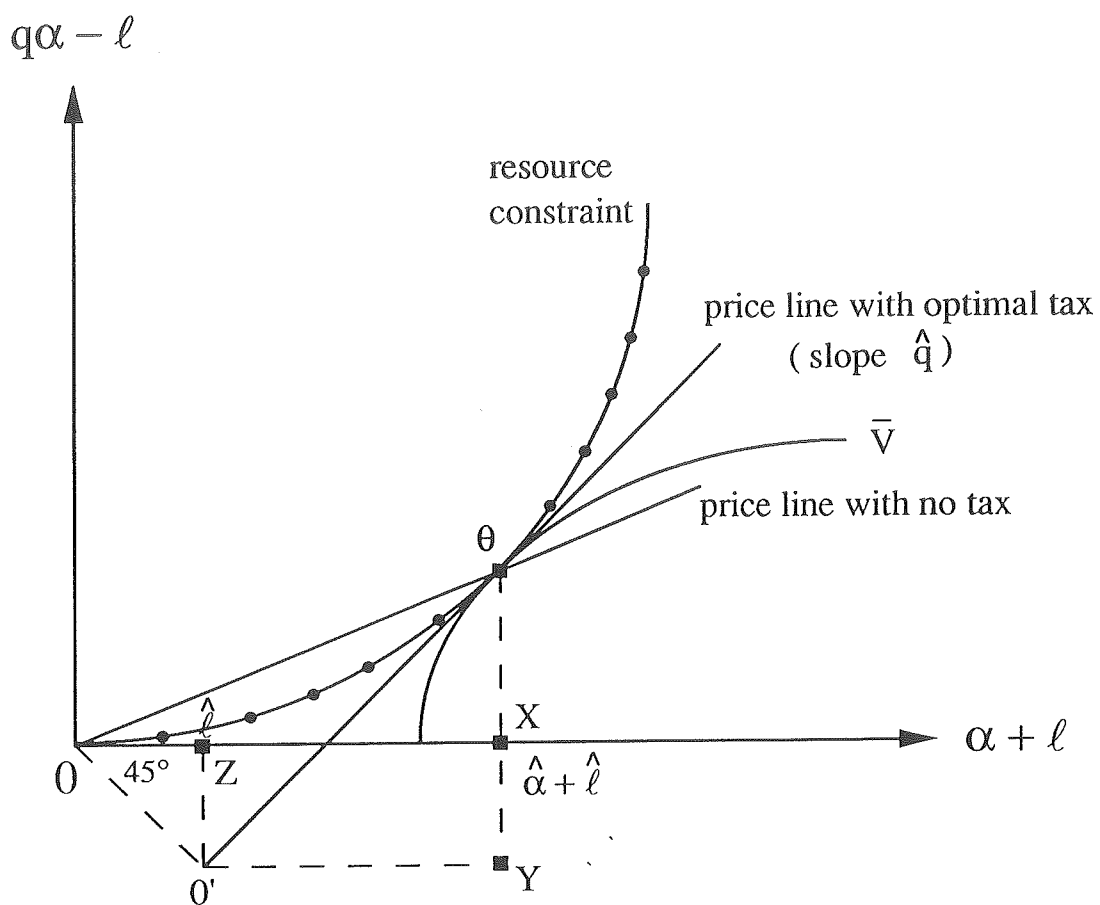


Figure 9: Convex indifference curves.

Note: $\hat{q} = \frac{\theta Y}{O'Y}, \hat{t} = \frac{\hat{\ell}}{\hat{\alpha}} = \frac{OZ}{O'Y} = \frac{O'Z}{O'Y} = \frac{XY}{O'Y},$

$$\hat{q} = \frac{\hat{p} + \hat{t}}{1 - \hat{p}}$$

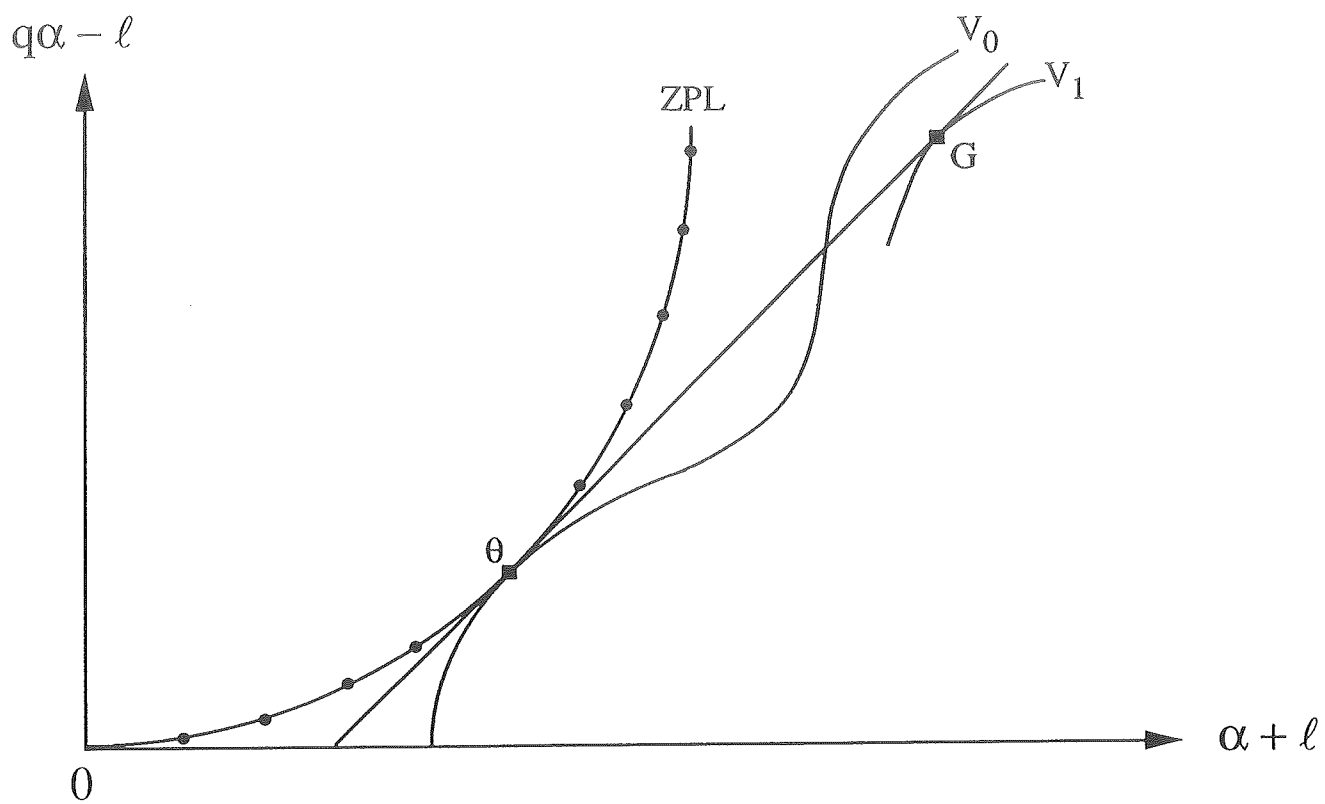


Figure 10: Non-convex indifference curves

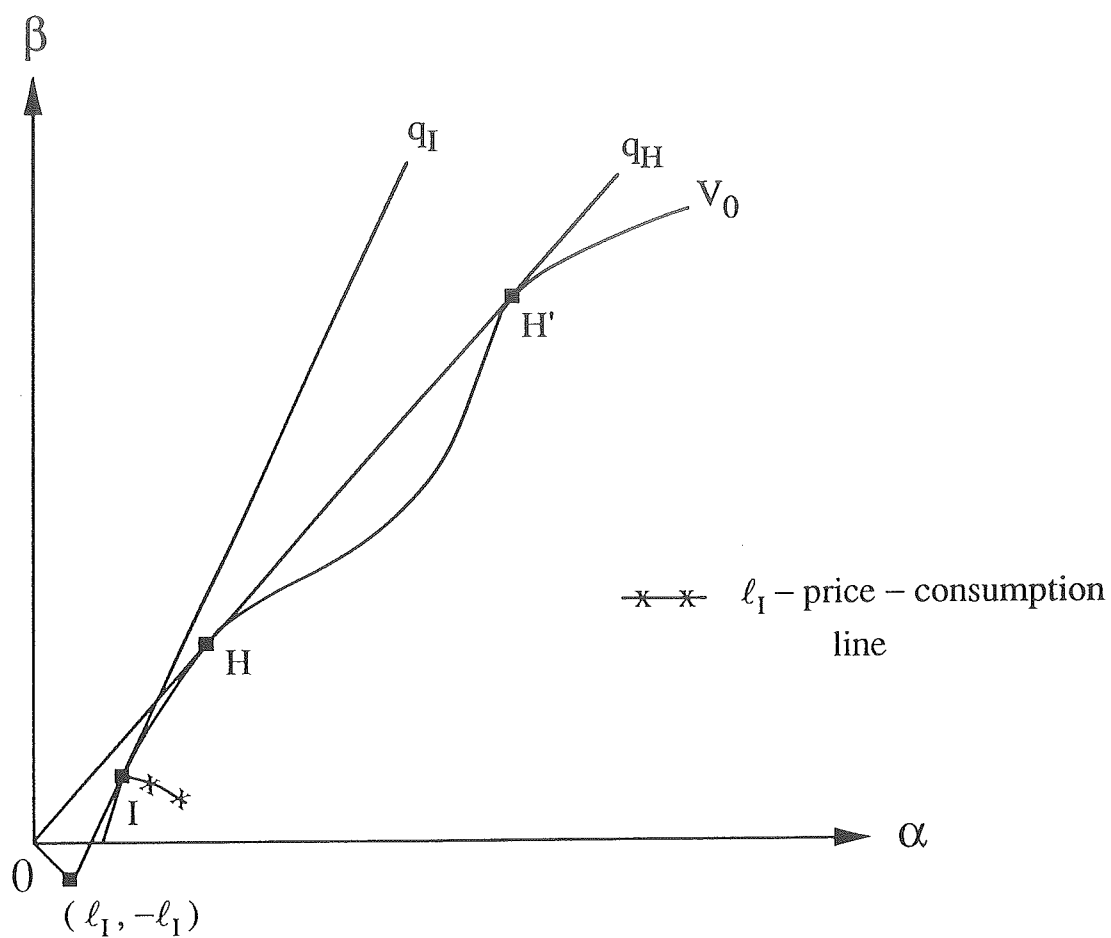


Figure 12: A small tax applied to a positive profit price equilibrium is always welfare-improving

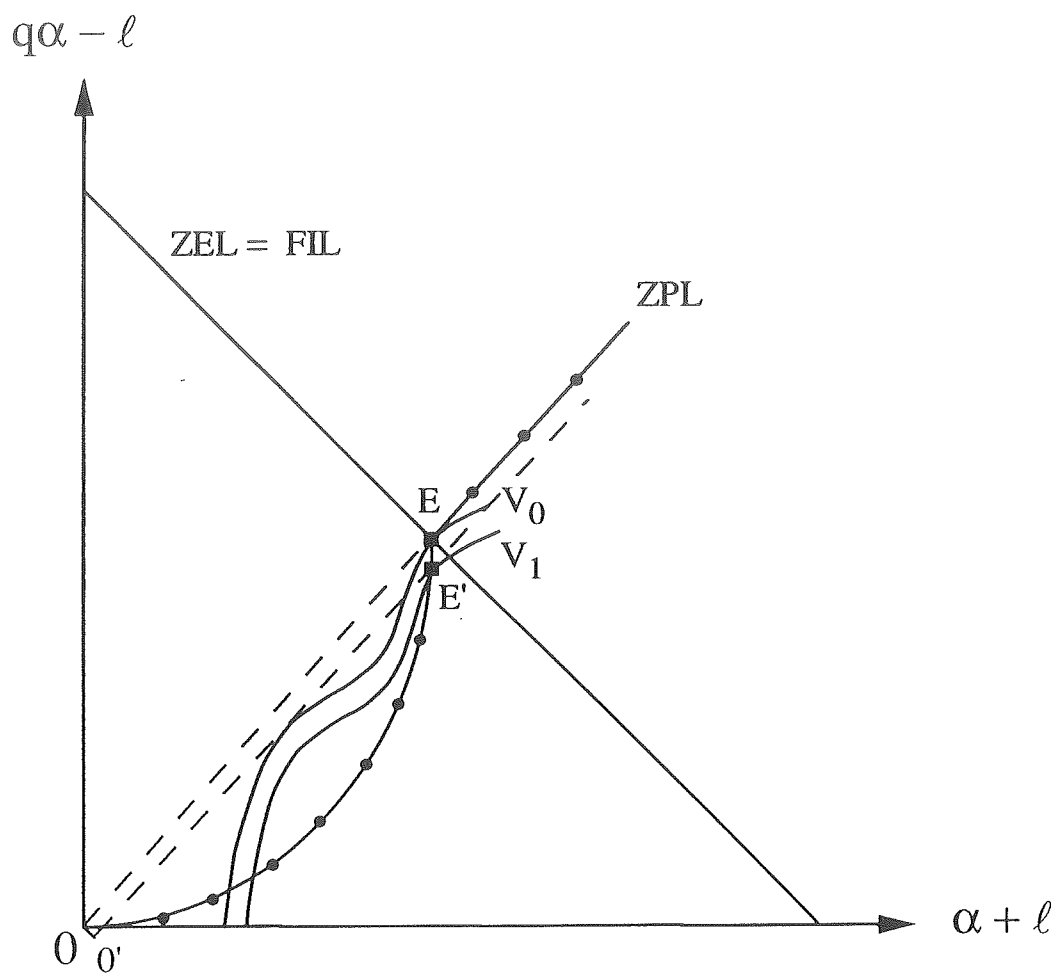


Figure 13: In case I, a small tax unambiguously increases welfare

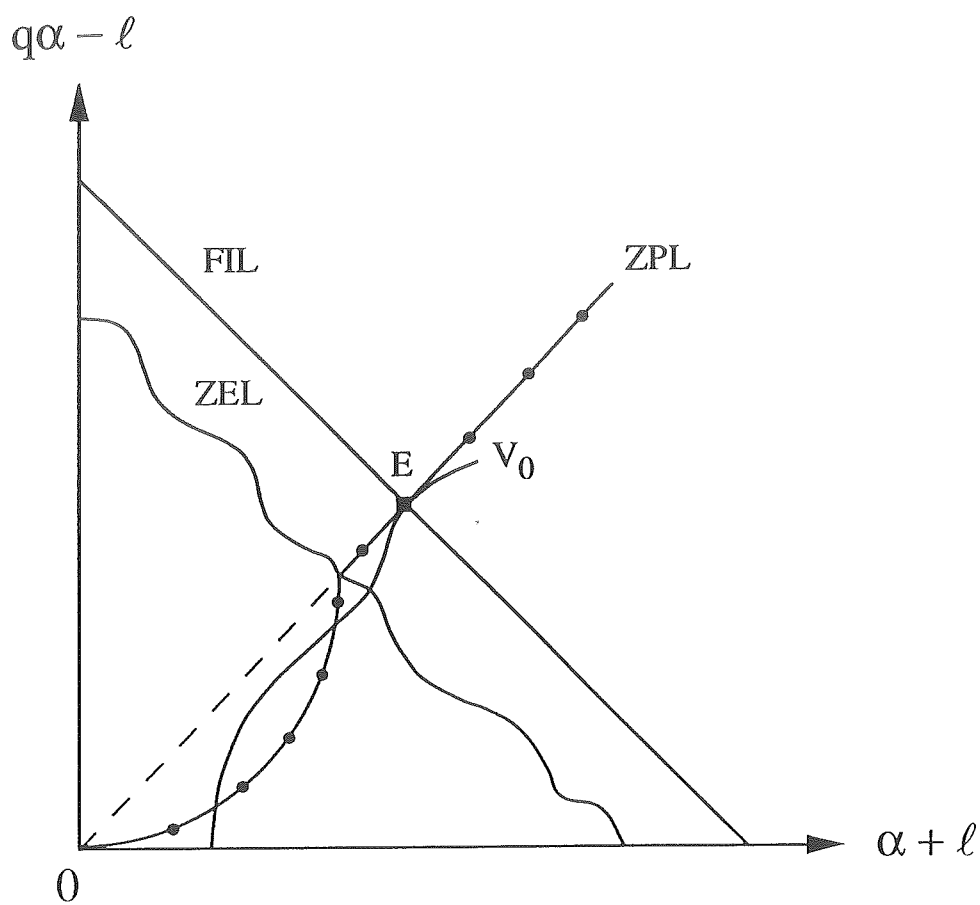


Figure 14: In cases III, IV, and V, a small tax unambiguously decreases welfare

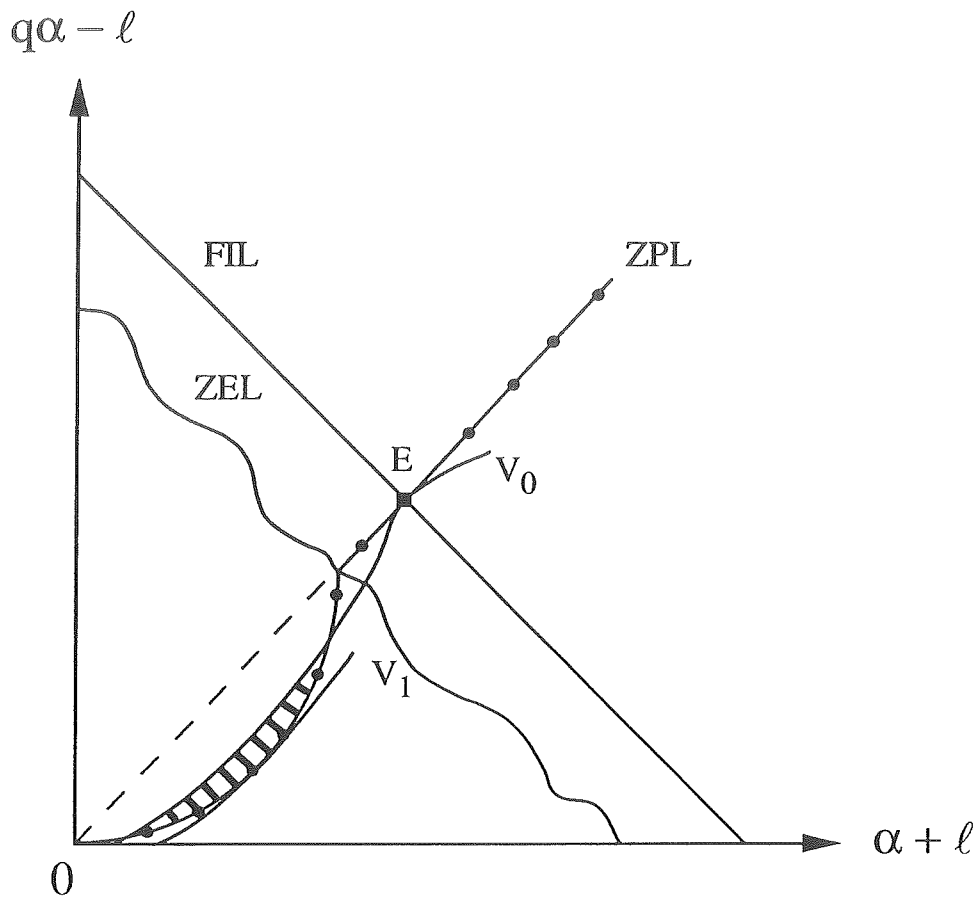


Figure 15: With a zero profit price equilibrium, there may not be a large, welfare-improving tax.