

PARETO-OPTIMAL INSURANCE POLICIES IN THE MODELS WITH A PREMIUM BASED ON THE ACTUARIAL VALUE

A. Y. Golubin

ABSTRACT

This article analyzes the problem of designing Pareto-optimal insurance policies when both the insurer and the insured are risk averse and the premium is calculated as a function of the actuarial value of the insurer's risk. Two models are considered: in the first, the set of admissible policies is constrained by a given size of the premium; in the second, the premium size is not constrained so that it varies with the actuarial value of a policy chosen by the agents. For both cases a characterization of the Pareto-optimal policies is derived. The corresponding optimality equations for the Pareto-optimal policies are obtained and compared with the results on the classical risk exchange model.

INTRODUCTION

Pareto-optimal risk exchanges in insurance have been studied in a range of papers, beginning with the works by Borch (1960) and Arrow (1971). Borch studied Pareto-optimal treaties between (re)insurance companies having initial stochastic losses. His theorem characterizing Pareto-optimal risk exchanges in a reinsurance market was extended to a constrained case in Gerber (1978). Borch's model, which is a special kind of n -person cooperative game, has been developed by Lemaire (1979, 1990) and Baton and Lemaire (1981) where such characteristics of the reinsurance market as the Pareto-optimal payoffs, core, bargaining set, and value of the game were investigated. Aase (2002) considered the model in a more general setting, including the competitive equilibrium notion and risk allocation problems in incomplete financial markets. In Arrow's formulation, the risk exchange is between two agents, the insurer and the insured, and the claim is supposed nonnegative. It is proved that if the insurer is risk neutral then the Pareto-optimal policy is one with a deductible and full coverage above the deductible. Raviv (1979) examined a model where both the agents are risk averse and administrative costs are an increasing and convex function of the covered loss. In this context, the Pareto-optimal policy is a contract with a deductible and a coinsurance clause above the deductible.

A. Y. Golubin is at the department of operations research, Moscow Institute of Electronics and Mathematics, Moscow, Russia. The author can be contacted via e-mail: io@miem.edu.ru.

Optimality of deductible policies was studied by Murray (1971) and then with more detailed and extensive analysis by Schlesinger (1981). Blazenko (1985) presented various conditions on costs in the model under which a policy with a deductible, coinsurance, or premium rebate is optimal. In Golubin (2002), a problem of parametrical optimization of insurance policies was studied under various assumptions on the summary risk distribution. Holmstrom (1979) analyzed the moral hazard problem and proved that this results in optimality of deductible policies.

In this article we study the problem of designing Pareto-optimal insurance policies (or, in other words, indemnity functions) in the situation where both the insurer and the insured are assumed to be risk averse, whereas the premium paid by the insured is based on the actuarial value of the insurer's risk. This way of premium calculation is believed to be quite natural from a view point of actuarial practice. In particular, it includes the expected value principle. The suggested formulation is not in the frame of the classical model of risk exchange because now the premium enters the payoffs, i.e., agents' expected utilities, as a functional of indemnity function, not as a constant or independent variable. The models in the present article are close to that by Raviv (1979) in which the administrative costs are set to be zero. The substantial difference is that we suppose the premium to be a given function $R(EI(X))$ of the actuarial value of the indemnity payment $EI(X)$. Thus, the optimization is taken in insurance policy $I(\cdot)$ only, whereas in Raviv's model the optimization is taken jointly in $I(\cdot)$ and premium size P . The analysis uses the Neyman–Pearson lemma instead of optimal control theory techniques, which allows for including into consideration the loss distribution of a general form instead of that having a positive density.

Two models of the contract are investigated. In the first model the set of admissible policies is constrained by a fixed size of the premium. In the second, this constraint is removed so that the premium varies with the actuarial value which, in turn, depends on an insurance policy chosen.

For both models, it is found that the Pareto-optimal policy belongs to one of the two kinds, with either a deductible or an upper limit of full coverage, and coinsurance above the deductible/upper limit. It is shown that the policy's kind is determined by a comparison of the value of a Pareto multiplier and the "critical" weight. The deductible Pareto-optimal policies appear under all sufficiently small values of the Pareto multiplier. Thus, the premium calculation based on the actuarial value is another "source" of deductible policies in addition to the settings with administrative costs or moral hazard.

The derived optimality equations are in terms of the Pareto multiplier and the second derivatives of the agents' utility functions, in distinction to the known optimality equation involving the risk-aversion functions (Raviv, 1979, Theorem 1). As is shown, solving the derived equations leads to Pareto-optimal policies different from those in the classical model.

The article is organized as follows. The next section describes a model of insurance contract under investigation and a procedure for finding Pareto-optimal policies. The section "Model with Constrained Premium" presents the results (in the form of necessary and sufficient optimality conditions) on constructing all Pareto-optimal policies in the model with constrained premium. In the section "Examples," these results are

illustrated by several examples of the agents' utility functions. The section "Model with Unconstrained Premium" presents a characterization of Pareto-optimal policies, where the case of a linear premium size function is discussed in detail.

PROBLEM FORMULATION

In this article the insurance market is treated as consisting of two agents, an insurance company and an individual facing a risk of loss (the insurer and the insured). The loss of the insured is a nonnegative stochastic value X with a distribution function $F(x) = P\{X \leq x\}$. We denote by $\text{supp } F$ the support of $F(x)$ (i.e., the least closed set $S \subseteq R$ such that $P\{X \in S\} = 1$) and by $T \stackrel{\text{def}}{=} \sup\{\text{supp } F\} \leq \infty$. Thus X may be a discrete stochastic value with a discrete set $\text{supp } F$ or, for instance, may have an exponential distribution with $\text{supp } F = [0, \infty)$.

An insurance policy or, in other words, an indemnity function, is identified with a measurable function $I(x)$ on $[0, \infty)$ satisfying the standard constraint $0 \leq I(x) \leq x$ (which implies in particular that $I(0) = 0$). At the end of the insurance period, the stochastic sum $I(X)$ will be paid by the insurer, and the residual loss $X - I(X)$ will be covered by the insured himself. Two policies, $I_1(x)$ and $I_2(x)$ are meant to be identical if they coincide for $x \in \text{supp } F$, that is, $P\{I_1(X) = I_2(X)\} = 1$. Provision of insurance does not involve any administrative expenses for the insurer in addition to his indemnity payment $I(X)$. The premium, i.e., the price paid by the insured, is denoted by P and assumed to be a function of the indemnity actuarial value, $P = R(E I(X))$, where $R(y)$ is a given continuous increasing function on $[0, E X]$ such that $R(0) = 0$. In the particular case of a linear function $R(y) = (1 + \alpha)y$ with $\alpha > 0$, we have the expected value principle of premium calculation that is widely used in insurance literature: $P = (1 + \alpha)E I(X)$ where α is the loading coefficient.

We suppose that the insurer and insured have, correspondingly, utility functions $u_0(x)$ and $u_1(x)$ such that they are twice continuously differentiable, $u'_i(x) > 0$ and $u''_i(x) < 0$, $i = 0, 1$,—thereby, both agents are supposed to be risk averse. The goal of each agent is to maximize the expected utility of his/her final capital by means of the choice of an insurance policy $I(x)$, the same for both agents. The insurer's expected utility is $J_0[I] \equiv E u_0(w_0 + P - I(X))$ and the insured's expected utility is $J_1[I] \equiv E u_1(w_1 - P - X + I(X))$, where w_i are initial nonstochastic capitals.

Finding Pareto-Optimal Solutions

For our case, a policy $\hat{I}$ is called *Pareto-optimal* if there is no other admissible I such that $J_i[I] \geq J_i[\hat{I}]$ for $i = 0, 1$ and at least one of the inequalities is strict. The proposition below presents a known procedure generating a one-parameter family of Pareto-optimal solutions by maximizing a weighted sum of the criteria. The result can be found in Gerber (1978, section 3) (see also Pardalos, Siskos, and Zopounidis, 1995).

Proposition 1:

(1) Let $\rho \in (0, 1)$. If $\hat{I}$ is a solution to the problem

$$\max_I \rho J_0[I] + (1 - \rho) J_1[I]$$

then $\hat{I}$ is Pareto-optimal.

- (2) Let $J_0[I]$ and $J_1[I]$ be concave functionals. If $\hat{I}$ is Pareto-optimal then there exists a $\rho \in [0, 1]$ such that $\hat{I}$ is a solution to the problem above.

For convenience, we rewrite the goal functional above as

$$J[I] \equiv \delta E u_0(w_0 + P - I(X)) + E u_1(w_1 - P - X + I(X)), \quad (1)$$

where $\delta \stackrel{\text{def}}{=} \rho/(1 - \rho)$ will be referred to as an insurer's relative weight. Evidently, $0 < \delta < \infty$ and, for example, the value $\delta = 1$ means that $\rho = 1 - \rho = 1/2$ —both agents are equally “significant” to each other in the contract. When δ runs over the interval $(0, \infty)$, solving (1) yields the family of Pareto-optimal policies $\{\hat{I}_\delta\}$ as was indicated in Proposition 1. If the functional $J[I]$ in (1) is concave in I , this procedure gives all the Pareto-optimal policies (except, possibly, those related to degenerated cases $\rho = 0 (=1)$).

There are several ways to single out a “best” Pareto-optimal policy. One can solve (1) under a concrete value of δ given exogenously, say, by “experts.” Such a procedure assigning the weights to agents' utilities with a use of regression methods and the social judgment theory is described in Pardalos, Siskos, and Zopounidis (1995). Two other optimality concepts introduced by Nash (1950), and by Kalai and Smorodinsky (1975) are based on the notion of a bargaining game or, in terminology of Lemaire (1979), a cooperative game without transferable utilities (see Borch 1960, 1990; and Lemaire 1979, 1990 for the definition and applications in insurance). Both concepts assume that the agents act rationally, that is, an optimal policy I should be Pareto-optimal and individually rational. The latter means $J_i[I] \geq J_i[0]$ (the agent will not agree to lessen his initial utility), where the initial point $(J_0[0], J_1[0]) = (u_0(w_0), E u_1(w_1 - X))$.

Problems of applying the above-mentioned optimality concepts are left beyond the scope of the article. In the sequel we focus only on obtaining a solution to (1), i.e., the Pareto-optimal policy $\hat{I}(x) = \hat{I}_\delta(x)$, for each fixed $\delta \in (0, \infty)$. This is justified by the fact that the main difficulty in applying the solution concepts seems to be the constructing of the set $\{\hat{I}_\delta\}_{\delta > 0}$ of Pareto-optimal policies.

MODEL WITH CONSTRAINED PREMIUM

Consider a situation where the premium paid by the insured for the insurance coverage must be a prescribed constant P , so that the insurer's admissible policies $I(\cdot)$ are those providing this premium size, i.e., satisfying the equation $R(E I(X)) = P$. Translating this restriction into $E I(X) = M$, where $M = R^{-1}(P)$ and $R^{-1}(\cdot)$ denotes the inverse function, we find that Equation (1) reduces in our case to maximization of

$$\begin{aligned} J[I] &\equiv \delta E u_0(w_0 + P - I(X)) + E u_1(w_1 - P - X + I(X)) \\ &\text{subject to } E I(X) = M, \quad 0 \leq I(x) \leq x. \end{aligned} \quad (2)$$

To exclude the degenerated cases in which the only admissible policy is $I(x) \equiv 0$ or $I(x) \equiv x$, the given constant P is assumed to be $R(0) < P < R(E X)$, that is, $0 < M < E X$. Both the expected utilities in (2) are defined on the convex set of admissible policies and concave in I due to concavity of $u_i(x)$. Then, by Proposition 1, solving (2) for all

fixed $\delta > 0$ gives the set $\{\hat{I}_\delta(x)\}_{\delta>0}$ of all Pareto-optimal policies (from here on we will omit the subscript δ for convenience).

A structure of the Pareto-optimal policy is specified by the following theorem that provides necessary and sufficient conditions for optimality in (2). Denote by $A = w_0 + P - \hat{I}(x)$ and $B = w_1 - P - x + \hat{I}(x)$ the final capitals of, respectively, the insurer and insured under a policy $\hat{I}$ if the loss $X = x$.

Theorem 1: *An admissible policy $\hat{I}(x)$ solves Equation (2) if and only if it takes one of the two forms*

$$\begin{cases} \hat{I}(x) = x, & \text{if } x \leq \hat{x} \\ 0 < \hat{I}(x) < x, & \text{if } x > \hat{x} \end{cases} \quad \text{for some } 0 \leq \hat{x} < T, \text{ or} \quad (3)$$

$$\begin{cases} \hat{I}(x) = 0, & \text{if } x \leq \hat{x} \\ 0 < \hat{I}(x) < x, & \text{if } x > \hat{x} \end{cases} \quad \text{for some } 0 \leq \hat{x} < T, \quad (4)$$

$$\text{and } \delta u'_0(A) - u'_1(B) \equiv \text{const for } x \in [\hat{x}, T]. \quad (5)$$

Moreover, optimality of the policy $\hat{I}$ having form (3) or (4) is equivalent to

$$\hat{I}'(x) = \frac{u''_1(B)}{\delta u''_0(A) + u''_1(B)} \quad (6)$$

on $[\hat{x}, T]$ with initial condition, respectively, either $\hat{I}(\hat{x}) = \hat{x}$ or $\hat{I}(\hat{x}) = 0$.

Remark 1: To completely define an indemnity $I(X)$, it suffices to define the function $I(x)$ on $\text{supp } F$ only. Therefore, equalities (3)–(5) in Theorem 1 are understood as satisfied on the indicated intervals up to a set of zero F -measure. On the other hand, condition (6) means the following: (i) if an admissible $\hat{I}(x)$ defined on $[0, T]$ is of form (3) or (4) and solves (6) then $\hat{I}$ is optimal; (ii) if $\hat{I}$ is optimal then there is a function $\bar{I}(x)$ on $[0, T]$ satisfying (3) (or (4)) and (6) such that $\bar{I}(x) \equiv \hat{I}(x)$ on $\text{supp } F$. Briefly, this means that an optimal policy $\hat{I}(x)$ can be determined via (3), (4), and (6) as a function on $[0, T]$, regardless of the shape of $\text{supp } F$.

Proof: Let $\hat{I}(x)$ denote a maximizer in Equation (2). Fix any admissible policy $I(x)$ and consider the policy $\lambda \hat{I}(x) + (1 - \lambda)I(x)$ where $\lambda \in [0, 1]$ is a parameter. It is seen that such a policy satisfies the constraints in (2). By optimality of $\hat{I}(x)$, the point $\lambda = 1$ maximizes $J(\lambda) = J[\lambda \hat{I} + (1 - \lambda)I]$ on the interval $[0, 1]$, hence the derivative at this point must be nonnegative

$$J'(\lambda)|_{\lambda=1} \geq 0$$

(note that the fulfillment of the inequality for any admissible policy $I(x)$ is necessary and sufficient for optimality in (2) because $J[\lambda \hat{I} + (1 - \lambda)I]$ is concave in λ due to the concavity of $u_0(x)$ and $u_1(x)$). After differentiating we have

$$E\{\delta u'_0(w_0 + P - \hat{I}(X)) - u'_1(w_1 - P - X + \hat{I}(X))[I(X) - \hat{I}(X)]\} \geq 0$$

or, equivalently, $\int_0^T \hat{U}(x)I(x) dF(x) \geq \int_0^T \hat{U}(x)\hat{I}(x) dF(x)$, where

$$\hat{U}(x) = \delta u'_0(w_0 + P - \hat{I}(x)) - u'_1(w_1 - P - x + \hat{I}(x)). \quad (7)$$

Thus, $\hat{I}(x)$ is a solution to (2) if and only if it minimizes the integral

$$\int_0^T \hat{U}(x)I(x) dF(x) \text{ subject to } \int_0^T I(x) dF(x) = M, \quad 0 \leq I(x) \leq x. \quad (8)$$

A characterization of solutions to (8) is given by a result known as the generalized Neyman–Pearson lemma (see Lehmann, 1959 or Karlin and Studden, 1966). According to it, an admissible policy $\hat{I}(x)$ is optimal in (8) if and only if there exists a constant c such that

$$\hat{I}(x) = \begin{cases} 0, & \text{if } \hat{U}(x) > c \\ x, & \text{if } \hat{U}(x) < c \end{cases} \text{ up to a set of zero } F\text{-measure.} \quad (9)$$

Suppose first that $\hat{U}(0) < c$. Then, as x increases from $x = 0$, we have by (9) that $\hat{I}(x)$ remains equal to x up to a point $\hat{x}$ at which $\hat{U}(x)$ first becomes equal to c —note that $\hat{U}(x)$ is increasing (see Equation (7)) with $\hat{I}(x) = x$. On the remaining interval $[\hat{x}, T]$, $\hat{U}(x)$ can neither decrease ($\hat{I}(x) = x$ for $\hat{U}(x) < c$, so $\hat{U}(x)$ would have increased for such x) nor increase (by (9), we would have $\hat{I}(x) = 0$ for $\hat{U}(x) > c$ and, according to (7), $\hat{U}(x)$ must decrease). Thus, the function $\hat{U}(x)$ is increasing on the interval $[0, \hat{x}]$, where $\hat{I}(x) = x$, and is identically equal to c on the interval $[\hat{x}, T]$ (more exactly, on $[\hat{x}, T] \cap \text{supp } F$). The latter is not empty since otherwise $\hat{I}(x) \equiv x$ with respect to F -measure and the constraint $E\hat{I}(X) = M$ with $0 < M < EX$ is violated.

Suppose that $\hat{U}(0) > c$. Similar to the reasonings above, we have that for some $0 < \hat{x} \leq T$ the function $\hat{U}(x)$ is decreasing on $[0, \hat{x}]$ with $\hat{I}(x) = 0$ there, and $\hat{U}(x) = c$, that is, $\delta u'_0(A) - u'_1(B) = c$ on the interval (nonempty) $[\hat{x}, T]$. Finally, if $\hat{U}(0) = c$ then, evidently, $\hat{U}(x) \equiv c$ and (5) is satisfied on the whole interval $[\hat{x}, T] = [0, T]$.

In order to prove (6) (see Remark 1), suppose first that an admissible $\hat{I}(x)$ defined on $[0, T]$ and having form (3) or (4) solves (6) on $[\hat{x}, T]$. Note that formal differentiating the equation $\hat{U}(x) = c (= \hat{U}(\hat{x}))$ on $[\hat{x}, T]$, where an expression for $\hat{U}(x)$ is given in (7), yields (6). The converse is also true, therefore $\hat{I}(x)$ satisfies (5) and, hence, is optimal.

Let now $\hat{I}(x)$ be optimal, i.e., satisfying (3)–(5) on $\text{supp } F$. Denote by $\bar{I}(x)$ a unique solution to (6) on $[\hat{x}, T]$ with the corresponding initial condition (existence and uniqueness of a solution to the equation of kind (6) were established, e.g., in Wyler (1990)). As was noted, (5) holds on $[\hat{x}, T]$ with $\bar{I}(x)$. Furthermore, (5) can be rewritten as $\delta u'_0(w_0 + P - \hat{I}(x)) - \hat{U}(\hat{x}) = u'_1(w_1 - P - x + \hat{I}(x))$ so that

$$\bar{I}(x) = f^{-1}(w_1 - P - x), \quad (10)$$

where f^{-1} is the inverse of the function $f(z) = (u'_1)^{-1}(\delta u'_0(w_0 + P - z) - \hat{U}(\hat{x})) - z$. Note that f^{-1} exists because $u'_0(\cdot)$ and $u'_1(\cdot)$ are decreasing functions. In view of (10), by uniqueness of solution we have $\bar{I}(x) \equiv \hat{I}(x)$ on $\text{supp } F$. Q.E.D.

Theorem 1 states that the optimal policy in (2) (i.e., the Pareto-optimal policy, according to Proposition 1) is either of form (3) or of form (4). Following the notation in Raviv (1979), in the sequel we will refer to these kinds of $\hat{I}(x)$ as, correspondingly, a *policy with an upper limit*, where $\hat{x}$ is the largest loss fully covered by the insurer, and a *policy with a deductible* in which $\hat{x}$ means the largest loss not covered by the insurer. From (6) it is seen that $0 < \hat{I}'(x) < 1$ for $x > \hat{x}$, therefore the Pareto-optimal contract includes a coinsurance for the loss greater than $\hat{x}$.

The derived optimality condition (5), $\delta u'_0(A) - u'_1(B) = \text{const}$ for $x \in [\hat{x}, T] \cap \text{supp } F$, looks similar to the Pareto optimality condition in Borch's model (Borch, 1960) of a reinsurance market: $\delta u'_0(A) - u'_1(B) = 0$ for $x \in \text{supp } F$. The constant in the right-hand side of (5), given by the Neyman–Pearson lemma (see Equation (9)) and generally not equal to zero, reflects the presence of moment constraint $E I(X) = M$ defining the set of admissible policies.

Unlike Raviv's model with a fixed premium (1979, section II), where the Pareto-optimal policy is determined by the optimality equation using the agents' risk-aversion functions, in our case $\hat{I}(x)$ is determined by a pair of conditions: Equation (6) (or its equivalent (5)) and equation $E \hat{I}(X) = M$.

Now we focus on a question still remaining open: what condition on the model parameters defines a kind of the Pareto-optimal risk exchange $\hat{I}(x)$? Let $I_0(\delta, x)$ be the function solving (6) on $[0, T]$ with initial condition $I_0(\delta, 0) = 0$ under a fixed δ . Define a "critical weight" δ^* as a unique solution¹ of the following equation with respect to δ

$$E I_0(\delta, X) = M. \quad (11)$$

So $I_0(\delta^*, x)$ is an admissible policy with neither deductible nor upper limit of full coverage.

Proposition 2: *The Pareto-optimal $\hat{I}(x)$ is a policy with an upper limit (with a deductible) $\hat{x} > 0$ if and only if the insurer's relative weight $\delta > (<) \delta^*$. In both cases the value $\hat{x}$ is a unique solution of the equation $E \hat{I}(X) = M$.*

The proof is given in the Appendix.

Remark 2: Let the relative weight δ infinitely increase then, as the insurer's and insured's weights (see the section "Problem Formulation") are $\rho = \delta/(1 + \delta) \rightarrow 1$ and $1 - \rho = 1/(1 + \delta) \rightarrow 0$, Equation (2) turns out to be the maximization of insurer's utility $E u_0(w_0 + P - I(X))$ only. As $\delta \rightarrow \infty$, the right-hand side of (6) tends to zero and from Proposition 2 we obtain that $\hat{I}(x)$ becomes a stop-loss policy with no

¹ The uniqueness of a solution to (11) follows from the fact that $I_0(\delta, x)$ is decreasing in δ , which is in turn independently established in the first part of Proposition 2's proof (see the Appendix).

coinsurance, $\hat{I}(x) = \min\{x, \hat{x}\}$. Oppositely, if $\delta \rightarrow 0$ then (2) converts into the problem of maximization of the insured's utility $E u_1(w_1 - P - X + I(X))$ and in this limiting case we clearly obtain a strict deductible policy (without coinsurance), $\hat{I}(x) = (x - \hat{x})_+ \stackrel{\text{def}}{=} \max\{0, x - \hat{x}\}$. A policy of the same kind arises if the insurer is assumed to be risk neutral. Indeed, $u_0''(x) \equiv 0$ and from (6) we have $\hat{I}'(x) \equiv 1$ on $[\hat{x}, T]$ for any fixed δ . Hence the only Pareto-optimal policy is $\hat{I}(x) = (x - \hat{x})_+$. Note that the level $\hat{x} \in (0, T)$ in all these cases is uniquely defined by solving the equation $E \hat{I}(X) = M$.

In a general case of a given $0 < \delta < \infty$, the following way for computing the Pareto-optimal $\hat{I}(x)$ can be suggested, based on Theorem 1 and Proposition 2. First, obtain in terms of δ a solution $I_0(\delta, x)$ to (6) with initial condition $I_0(\delta, 0) = 0$. Second, find the "critical value" δ^* by solving the equation $E I_0(\delta, X) = M$. Third, compare the given δ and δ^* : if $\delta \geq \delta^*$ then obtain a policy $I_1(\hat{x}, x)$ with an upper limit $\hat{x} \geq 0$, that satisfies (6) on $[\hat{x}, T]$ and $I_1(\hat{x}, x)|_{x=\hat{x}} = \hat{x}$, with $\hat{x}$ being a parameter. Fourth, determine an optimal upper limit $\hat{x}$ as a unique solution of the equation $E I_1(\hat{x}, X) = M$. Analogously, if $\delta < \delta^*$ then obtain from (6) an expression for a policy $\hat{I}_2(\hat{x}, x)$ with a deductible $\hat{x} > 0$ and uniquely determine the value of $\hat{x}$ from $E I_2(\hat{x}, X) = M$. Remark that Equation (2) always has a unique solution $\hat{I}(x)$ (dependent on δ).

EXAMPLES

In this section we illustrate finding the Pareto-optimal policy via the treatment of optimality Equation (6). We consider three examples that involve exponential, quadratic, and logarithmic utility functions.

Example 1: Let the insurer and insured have, respectively, quadratic utility functions $u_i(x) = -\frac{1}{2}\beta_i x^2 + x$ with the domains $x < 1/\beta_i$, $i = 0, 1$.

In this case the right-hand side of (6) is constant and $\hat{I}'(x) = \theta$ on $[\hat{x}, T]$, where $\theta \stackrel{\text{def}}{=} \beta_1/(\delta\beta_0 + \beta_1)$. From (3) and (4) it follows that the Pareto-optimal policy with an upper limit must be of the form

$$\hat{I}(x) = I_1(\hat{x}, x) = x - (1 - \theta)(x - \hat{x})_+, \quad (12)$$

and the Pareto-optimal policy with a deductible must be of the form

$$\hat{I}(x) = I_2(\hat{x}, x) = \theta(x - \hat{x})_+. \quad (13)$$

To determine the kind of the Pareto-optimal $\hat{I}(x)$, note first that the policy $I_0(\delta, x)$ used in Proposition 2 is in this case $I_0(\delta, x) = \theta x$. From the equation $E I_0(\delta, X) = M$, we have that the critical weight

$$\delta^* = \frac{\beta_1}{\beta_0} \left(\frac{E X}{M} - 1 \right).$$

Thus, if the given $\delta \leq \delta^*$ then the Pareto-optimal $\hat{I}(x)$ is a policy $I_2(\hat{x}, x)$, where a deductible $\hat{x} \geq 0$ is determined by

$$E I_2(\hat{x}, X) = M \text{ or, equivalently, } \int_0^T (x - \hat{x})_+ dF(x) = M(1 + \delta\beta_0/\beta_1). \quad (14)$$

If $\delta > \delta^*$ then the Pareto-optimal $\hat{I}(x)$ is a policy $I_1(\hat{x}, x)$, where an upper limit $\hat{x} > 0$ is determined by

$$E I_1(\hat{x}, X) = M, \text{ that is, } \int_0^T (x - \hat{x})_+ dF(x) = (E X - M) \left(1 + \frac{\beta_1}{\delta\beta_0} \right). \quad (15)$$

Let, for instance, X be uniformly distributed on $[0, 1]$. Then for $\delta \leq \delta^* = (1/M - 2)\beta_1/(2\beta_0)$, from (14) we obtain that the deductible is $\hat{x} = 1 - \sqrt{2M[1 + \delta\beta_0/\beta_1]}$. For $\delta > \delta^*$, Equation (15) gives the upper limit $\hat{x} = 1 - \sqrt{(1 - 2M)[1 + \beta_1/(\delta\beta_0)]}$.

Example 2: Let both agents have exponential utility functions: $u_i(x) = (1 - \exp(-\beta_i x))/\beta_i$, $i = 0, 1$. As we will see, this case gives a nonlinear Pareto-optimal policy. Denote by $x(I)$ the inverse of the function $\hat{I}(x)$ on $[\hat{x}, T]$ and let $z(I) \equiv x(I) - I$. Then from (6) we obtain

$$\frac{dx}{dI} = a \frac{\beta_0}{\beta_1} \exp(\beta_0 I - \beta_1 z) + 1, \text{ hence, } \frac{dz}{dI} = a \frac{\beta_0}{\beta_1} \exp(\beta_0 I - \beta_1 z),$$

where the constant $a = \delta \exp[\beta_1(w_1 - P) - \beta_0(w_0 + P)]$. This separable differential equation can be rewritten as $\exp(\beta_1 z) dz = a \frac{\beta_0}{\beta_1} \exp(\beta_0 I) dI$ and, by integration, we have $z(I) = \beta_1^{-1} \ln[a \exp(\beta_0 I) + C]$, whence

$$x(I) = I + \beta_1^{-1} \ln[a \exp(\beta_0 I) + C]. \quad (16)$$

The constant C is defined by the corresponding initial condition in (6). For the case of an upper limit policy, the condition gives $x(I)|_{I=\hat{x}} = \hat{x}$ therefore (see Equation (16)) $C = 1 - a \exp(\beta_0 \hat{x})$. For a deductible policy, $x(I)|_{I=0} = \hat{x}$ and from (16) it follows that $C = \exp(\beta_1 \hat{x}) - a$.

In both cases, the Pareto-optimal policy on $[\hat{x}, T]$ is defined as

$$\hat{I}(x) = \psi(x), \quad (17)$$

where $\psi(x)$ denotes the inverse of the function $x(I)$ in (16). Since $x(I)$ is increasing and convex, $\hat{I}(x)$ is an increasing and concave function. The value of $\hat{I}(x)$ at a fixed $x \in [\hat{x}, T]$ can be easily obtained by numerical solving the equation $I + \beta_1^{-1} \ln[a \exp(\beta_0 I) + C] = x$ for I . Point out that expression (17) is quite different from a simple formula, $\hat{I}(x) = \beta_1/(\beta_0 + \beta_1)x$ on $[0, T]$, followed from Raviv's results (1979, pp. 87, 90) for his model with optimization in both $I(\cdot)$ and P in the case of exponential utilities (and without administrative costs as in our model).

The above-stated condition (11) for determining the critical weight δ^* , i.e., the kind of optimal $\hat{I}(x)$, converts into the following equation

$$\int_0^T \psi(x) dF(x) = M \text{ or } \int_0^{\psi(T)} y dF(x(y)) = M, \quad (18)$$

with $\hat{x} = 0$ and, hence, constant C in (16) equal to $1 - a$. (Note also that if $F(x)$ has a density $f(x)$ then $dF(x(y)) = f(x(y)) x'(y) dy$, and the integral in the latter equation becomes a Riemann integral.)

The equation $E \hat{I}(X) = M$ yielding the value of $\hat{x}$ (where, recall, the number $M = R^{-1}(P)$ is defined by a given premium size P) takes here one of the following forms. For the upper limit case ($\delta > \delta^*$):

$$\int_0^{\hat{x}} x dF(x) + \int_{\hat{x}}^{\psi(T)} y dF(x(y)) = M;$$

for the deductible case ($\delta < \delta^*$):

$$\int_0^{\psi(T)} y dF(x(y)) = M.$$

Example 3: Consider a situation where one of the agents has a quadratic utility function. Let, first, the insured's utility be quadratic, $u_1(x) = -\frac{1}{2}\beta_1 x^2 + x$, with the domain $x < 1/\beta_1$ which implies that the parameters satisfy $w_1 - P < 1/\beta_1$. Then Equation (6) converts into $I'(x) = [1 - u_0''(w_0 + P - I(x))\delta/\beta_1]^{-1}$ or, equivalently, $[1 - u_0''(a - I)\delta/\beta_1]dI = dx$, where $a = w_0 + P$. After integration we have $I + u_0'(a - I)\delta/\beta_1 = x + C$, $x \in [\hat{x}, T]$. Hence, the Pareto-optimal policy on this interval is

$$\hat{I}(x) = \psi(x), \quad (19)$$

where $\psi(x)$ is the inverse of the function $x(I) = I + u_0'(a - I)\delta/\beta_1 - C$.

The integration constant C is determined according to the kind of the Pareto-optimal policy. Upper limit case: $\hat{I}(\hat{x}) = \hat{x}$, hence (see Equation (19)), $C = u_0'(a - \hat{x})\delta/\beta_1$; deductible case: $\hat{I}(\hat{x}) = 0$, and from (19) it follows $C = u_0'(a)\delta/\beta_1 - \hat{x}$.

The relations establishing the kind of $\hat{I}(x)$ and the value of $\hat{x}$ have the same form as that in the previous example, with the function $\psi(x)$ defined as above.

In a particular case of logarithmic utility $u_0(x) = \ln x$, from (19) we get an analytical expression for the Pareto-optimal policy on $[\hat{x}, T]$,

$$\hat{I}(x) = \frac{1}{2} \left[x + C + a - \sqrt{(x + C - a)^2 + 4\delta/\beta_1} \right]. \quad (20)$$

Let the insurer now have a quadratic utility $u_0(x) = -\frac{1}{2}\beta_0 x^2 + x$ (with the domain $x < 1/\beta_0$), while $u_1(x)$ is an arbitrary utility function. If we denote $z(x) \equiv x - \hat{I}(x)$ for $x \in [\hat{x}, T]$, from (6) it follows that $[1 - u_1'(b - z)/(\delta\beta_0)]dz = dx$, where $b = w_1 - P$. Hence

$$z + u_1'(b - z)/(\delta\beta_0) = x + C. \quad (21)$$

By the definition of $z(x)$, the Pareto-optimal policy on the interval $[\hat{x}, T]$ is specified as

$$\hat{I}(x) = x - \psi(x), \quad (22)$$

where $\psi(x)$ is the inverse of the function $k(z) \stackrel{\text{def}}{=} z + u_1'(b - z)/(\delta\beta_0) - C$. In the case of an upper limit policy, the constant C is determined by $\hat{I}(\hat{x}) = \hat{x}$ which is equivalent to $z(\hat{x}) = 0$, whence (see Equation (21)) $C = u_1'(b)/(\delta\beta_0) - \hat{x}$. In the case of a deductible policy, the condition $\hat{I}(\hat{x}) = 0$ means $z(\hat{x}) = \hat{x}$, whence $C = u_1'(b - \hat{x})/(\delta\beta_0)$.

Particularly, if the insured's utility function is logarithmic, $u_1(x) = \ln x$, then (22) allows for deriving the optimal policy on $[\hat{x}, T]$ in an explicit form:

$$\hat{I}(x) = \frac{1}{2} \left[x - C - b + \sqrt{(x + C - b)^2 + 4/(\delta\beta_0)} \right]. \quad (23)$$

Notice that here $\hat{I}(x)$ is convex, whereas in a similar situation with a logarithmic insurer's utility function, $\hat{I}(x)$ (see Equation (20)) on the coinsurance interval $[\hat{x}, T]$ turns out to be a concave function.

MODEL WITH UNCONSTRAINED PREMIUM

We analyze here the Pareto-optimal insurance policy for the contract where the premium size P is not fixed and depends on a choice of insurance policy $I(x)$ through the actuarial value formula $P \equiv R(E I(X))$, where the function $R(y)$ is now additionally assumed to be differentiable on $[0, EX]$.

This situation seems more common in insurance than the previous model with constrained premium. As we will see, removing the constraint on the premium size makes the analysis more complicated. On the other hand, most of the results obtained in the previous section can be adjusted for this model.

In the considered case, Equation (2) of determining the Pareto-optimal policy converts into the following optimization problem: maximize

$$\begin{aligned} J[I] &\equiv \delta E u_0(w_0 + P - I(X)) + E u_1(w_1 - P - X + I(X)) \\ \text{subject to } &0 \leq I(x) \leq x, \end{aligned} \quad (24)$$

with $P \equiv R(E I(X))$. Let us first establish the existence of a solution to (24). This problem can clearly be formulated as maximization of the function

$$\phi(M) \stackrel{\text{def}}{=} \max\{J[I] \mid EI(X) = M, 0 \leq I(x) \leq x\}$$

over the interval $M \in [0, EX]$. At each M , $\phi(M)$ is the maximal value in Equation (2) with the same M in the constraint. As was shown in the section "Model With Constrained Premium," a maximizer $\hat{I}(x) = \hat{I}(M, x)$ in (2) is unique under any fixed $M \in [0, EX]$ (note that $\hat{I}(M, x) \equiv x \equiv 0$ if $M = EX (=0)$). When M decreases from the value EX , the policy $\hat{I}(M, x)$ is at first an upper limit policy with $\hat{x} = \hat{x}(M)$ continuously changing from T to 0 , where the point $0 = \hat{x}(M')$ corresponds to M' at which $\delta = \delta^*(M')$ (see Equation (11)). As M decreases from M' to 0 , $\hat{I}(M, x)$ becomes a deductible policy, the deductible $\hat{x}(M)$ continuously changes from 0 to T —the case $\hat{x} = T$ means that $\hat{I}(x) \equiv 0$ if $M = 0$. Since $\hat{I}(M, x)$ is a continuous function on $[0, EX] \times [0, T]$, $\phi(M) = E\{\delta u_0(w_0 + P - \hat{I}(M, X)) + u_1(w_1 - P - X + \hat{I}(M, X))\}$ is also continuous on the compact $[0, EX]$ by continuity of $u_i(x)$, $i = 0, 1$ and $R(y)$. Hence, $\max_{M \in [0, EX]} \phi(M)$ exists, therefore (24) has a solution. Thus, we have proved the following.

Lemma 1: *There exists an optimal policy $\hat{I}(x)$ in (24).*

The next theorem is an analogue of Theorem 1 and provides a characterization of the Pareto-optimal policy $\hat{I}(x)$ in a similar way (see also Remark 1). Denote the final agents' capitals under $X = x$ by $A(x) = w_0 + \hat{P} - \hat{I}(x)$ and $B(x) = w_1 - \hat{P} - x + \hat{I}(x)$, where now $\hat{P} = R(E\hat{I}(X))$.

Theorem 2: *If $\hat{I}(x)$ is a solution to problem (24) then for some $0 \leq \hat{x} \leq T$ it takes one of the two forms, (3) or (4), and*

$$\delta u'_0(A(x)) - u'_1(B(x)) \equiv \hat{C} \quad \text{for } x \in [\hat{x}, T] \quad (25)$$

where $\hat{C} = R'(E\hat{I}(X))E\{\delta u'_0(A(X)) - u'_1(B(X))\}$.

Moreover, (25) is equivalent to a pair of conditions (26)–(27):

$$\hat{I}'(x) = \frac{u''_1(B(x))}{\delta u''_0(A(x)) + u''_1(B(x))} \quad \text{for } x \in [\hat{x}, T] \quad (26)$$

with initial condition, respectively, either $\hat{I}(\hat{x}) = \hat{x}$ or $\hat{I}(\hat{x}) = 0$, and

$$\delta u'_0(A(\hat{x})) - u'_1(B(\hat{x})) = \hat{C}. \quad (27)$$

Proof: Similar to Theorem 1's proof, fix a maximizer $\hat{I}(x)$ in (24) and any admissible policy $I(x)$. By optimality of $\hat{I}$, the point $\lambda = 1$ maximizes the function $J(\lambda) \equiv J[\lambda\hat{I} + (1 - \lambda)I]$ on the interval $[0, 1]$, hence, the derivative $J'(1) \geq 0$. After differentiating, we have that this inequality becomes equivalent to that $\hat{I}(x)$ minimizes the integral

$$\int_0^T (\hat{U}(x) - \hat{C})I(x) dF(x) \quad \text{subject to } 0 \leq I(x) \leq x, \quad (28)$$

$$\text{where } \hat{U}(x) \stackrel{\text{def}}{=} \delta u'_0(A(x)) - u'_1(B(x)). \quad (29)$$

Applying the Neyman–Pearson lemma to (28) (in which, unlike (8) in Theorem 1, the moment constraint $E I(X) = M$ is absent), an admissible policy $\hat{I}(x)$ is optimal if and only if

$$\hat{I}(x) = \begin{cases} 0, & \text{if } \hat{U}(x) - \hat{C} > 0 \\ x, & \text{if } \hat{U}(x) - \hat{C} < 0. \end{cases}$$

In the same way as in Theorem 1, we examine the cases $\hat{U}(0) < \hat{C}$ and $\hat{U}(0) > \hat{C}$ and obtain, respectively, either $\hat{I}(x) = x$ or $\hat{I}(x) = 0$ on the first interval $[0, \hat{x}]$, whereas $\hat{U}(x) = \hat{C}$ on the interval (possibly empty) $[\hat{x}, T]$ in both cases. As in Theorem 1, the equalities are understood with respect to $\text{supp } F$, i.e., up to a set of zero F -measure. Note that emptiness of $[\hat{x}, T]$ means that either $\hat{U}(x) - \hat{C} < 0$ or $\hat{U}(x) - \hat{C} > 0$ on $[0, T]$, which, in turn, means that $\hat{I}(x) \equiv x$ or $\hat{I}(x) \equiv 0$, respectively.

To establish (26), we use a method similar to that used in proving (6). Let $\hat{I}(x)$ solve (26) on $[\hat{x}, T]$. Then it is easy to verify that $\hat{U}(x) \equiv \text{const}$ on $[\hat{x}, T]$. This equality coupled with $\hat{U}(\hat{x}) = \hat{C}$ (see Equation (27)) gives $\hat{U}(x) \equiv \hat{C}$ on $[\hat{x}, T]$.

Now let $\hat{I}(x)$ satisfy (25) on $[\hat{x}, T] \cap \text{supp } F$. Denote by $\bar{I}(x)$ a unique solution of (26) on $[\hat{x}, T]$, where the premium is fixed to be $\hat{P} = R(E \hat{I}(X))$. Then $\bar{I}$ solves (25) on $[\hat{x}, T]$ with the fixed $\hat{P}$. On the other hand, a unique solution to this equation (see Equation (10)) is $\bar{I}(x) = f^{-1}(w_1 - \hat{P} - x)$, where f^{-1} is the inverse of the function $f(z) = (u'_1)^{-1}(\delta u'_0(w_0 + \hat{P} - z) - \hat{C}) - z$. By the uniqueness, we have

$$\bar{I}(x) = \hat{I}(x) \text{ on } [\hat{x}, T] \cap \text{supp } F \quad (30)$$

and, as $\bar{I}(x) = \hat{I}(x)$ on $[0, \hat{x}] \cap \text{supp } F$, $R(E \bar{I}(X)) = R(E \hat{I}(X))$. The latter implies that $\bar{I}(x)$ solves (26) with $\hat{P} = R(E \bar{I}(X))$, which along with (30) completes the proof. Q.E.D.

Remark 3: Unlike Theorem 1, Theorem 2 assumes the degenerated policies $\hat{I}(x) \equiv x (\equiv 0)$. These cases mean that Equation (27) is unsolvable and, correspondingly, $\hat{U}(0) < \hat{C} (> \hat{C})$. The policy $\hat{I}(x) \equiv x$ (or $\equiv 0$) has evidently the sense of full coverage (or refusal of contract).

The derived optimality Equation (26) differs from Equation (6) in Theorem 1 by the presence of $\hat{P} = R(E \hat{I}(X))$ instead of the constant P . This makes (26) not an ordinary differential equation in the classic sense since the unknown function $\hat{I}(\cdot)$ enters its right-hand side also as an integrant in $\hat{P} = R(\int_0^T \hat{I}(t) dF(t))$. However, it is easy to see that this is equivalent to a pair of equations: the first is differential Equation (26) in which $\hat{P}$ is regarded as a parameter. After obtaining its parameterized solution $I(\hat{P}, x)$, the value of $\hat{P}$ is determined by the second equation $\hat{P} = R(E I(\hat{P}, X))$. In this connection, one can see that formulas (12)–(13), (19)–(20), (22)–(23) in Examples 1–3 actually provide expressions for such solutions under a fixed $P = \hat{P}$.

Let us examine two “degenerated” cases where either the insurer or insured is risk neutral. In the first case, $u''_0(x) \equiv 0$, hence the right-hand side of (26) becomes equal to 1 and, by Theorem 2, the Pareto-optimal $\hat{I}(x)$ is a deductible policy $\hat{I}(x) = (x - \hat{x})_+$ with some $\hat{x} \in [0, T]$. In the second case, from $u'_1(x) \equiv 0$ we similarly have that $\hat{I}(x) = 0$ on $[\hat{x}, T]$, hence $\hat{I}(x)$ is a stop-loss policy $\hat{I}(x) = \min\{x, \hat{x}\}$. One can note that these forms

of the Pareto-optimal policies are consistent with the results in Arrow (1971) and in Raviv's model with a fixed premium and zero administrative costs (1979, Theorem 1) if either the insurer or insured is supposed risk neutral.

It seems interesting to show how the classical result in Raviv (1979, Theorem 1) describing the form of the Pareto-optimal policy in the model with a fixed premium (and without administrative costs) follows from Theorem 2 as a special case. If $R(y) \equiv P$ then $R'(y) \equiv 0$, therefore $\hat{C} = 0$ and condition (25) becomes

$$\delta u'_0(w_0 + P - \hat{I}(x)) - u'_1(w_1 - P - x + \hat{I}(x)) = 0 \quad (31)$$

for $x \in [\hat{x}, T]$. Differentiating it with respect to x and then inserting an expression for δ obtained from (31) yield in this case the differential equation presented in Raviv's Theorem 1:

$$\hat{I}'(x) = \frac{r_1(B(x))}{r_0(A(x)) + r_1(B(x))},$$

where $r_i(x) = -u''_i(x)/u'_i(x)$, $i = 0, 1$ are the risk-aversion functions of the agents. Note that in a general case $R(\cdot) \neq \text{const}$ we get $\hat{C} \neq 0$, therefore optimality Equation (26) in Theorem 2 cannot be rewritten in terms of the risk-aversion functions.

An apriori characterization of the kind of $\hat{I}(x)$ in terms of the relative weight δ is given by Proposition 3 below. It is similar to Proposition 2 in the section "Model with Constrained Premium" but uses an extra condition on the premium size function $R(y)$. Suppose

$$R'(y) > 1, \quad (32)$$

that is, the premium $R(EI(X))$ increases faster than the actuarial value $EI(X)$. Let $I_0(\delta, x)$ be a policy that solves Equation (26) on $[0, T]$ with initial condition $I_0(\delta, 0) = 0$ (a policy with neither deductible nor upper limit), where $0 < \delta < \infty$ is considered as a parameter. Introduce an equation with respect to δ :

$$\delta u'_0(A(0)) - u'_1(B(0)) = \hat{C}, \quad (33)$$

where expressions for $A(\cdot)$, $B(\cdot)$, and $\hat{C}$ are the same as that in Theorem 2 but with the policy $I_0(\delta, x)$ instead of $\hat{I}(x)$.

Proposition 3: *Let (32) hold and δ^* be a unique solution to (33). Then the Pareto-optimal $\hat{I}(x)$ is a policy with an upper limit (with a deductible) $\hat{x} > 0$ if and only if $\delta > (<) \delta^*$.*

The proof is given in the Appendix.

Equation (33) plays the same role of determining a "critical weight" δ^* as Equation (11) in the model with a constrained premium. Like the results in the section "Model with Constrained Premium," Proposition 3 shows that if the insurer's relative weight δ is large enough ($\delta > \delta^*$) then $\hat{I}(x)$ is necessarily a policy with an upper limit, and $\hat{I}(x)$ must be a policy with a deductible for small δ ($\delta < \delta^*$). Finally, if $\delta = \delta^*$ then $\hat{I}(x)$ is the policy $I_0(\delta^*, x)$ with neither deductible nor upper limit. Once the kind of the Pareto-optimal policy is established, the form of $\hat{I}(x)$ is determined by Equation (26). The corresponding optimal level $\hat{x}$ (upper limit or deductible) in $\hat{I}(x) = \hat{I}(\hat{x}, x)$ is obtained

by condition (27) which plays the same part here as the equation $E I(\hat{x}, X) = M$ played in the previous model.

The Case of a Linear Premium Size Function R

Note first that Theorem 2 gives necessary optimality conditions only, whereas for the previous model with constrained premium Theorem 1 establishes the necessary and sufficient conditions for optimality. This is explained by the fact that the functional $J[I]$ in (24) may not be concave. However, in a particular case of a linear $R(y)$, it becomes concave by concavity of utility functions u_i and Theorem 2 provides then necessary and sufficient conditions for optimality in (24). Furthermore, the next result shows that under each fixed $\delta > 0$ a solution to (24) is unique. Let $T_0 \stackrel{\text{def}}{=} \min\{\text{supp } F\}$ then the equality $T_0 = 0$ means that the possible values of X are not separated from 0.

Corollary 1: *Let $R(y) = ly$ with $l > 0$, and $T_0 = 0$. Then (24) has a unique solution.*

Proof: To establish the solution uniqueness, it is sufficient to prove the strict concavity of $J[I]$ (recall that a solution to (24) exists by Lemma 1). Let $\lambda \in (0, 1)$ and $I_k(x)$, $k = 1, 2$ be any policies such that $I_1 \neq I_2$. Then $J[\lambda I_1 + (1 - \lambda)I_2] = E\{\delta u_0(w_0 + [P - I(X)]_\lambda - X)\} + u_1(w_1 - [P - I(X)]_\lambda - X)$, where we denote $[P - I(X)]_\lambda = \lambda(P_1 - I_1(X)) + (1 - \lambda)(P_2 - I_2(X))$ and $P_k = l E I_k(X)$, $k = 1, 2$. Note that $P_1 - I_1(\cdot) \neq P_2 - I_2(\cdot)$ since otherwise we would have either $P_1 = P_2$ and $I_1(\cdot) = I_2(\cdot)$, or $P_1 \neq P_2$ and $I_1(\cdot) = I_2(\cdot) + (P_1 - P_2)$ which leads to violating the condition $I_k(0) = 0$, $k = 1, 2$. Hence, by the strict concavity of u_0 and u_1 ,

$$J[\lambda I_1 + (1 - \lambda)I_2] > \lambda J[I_1] + (1 - \lambda)J[I_2]. \quad \text{Q.E.D.}$$

Remark 4: Under the linear $R(y) = ly$, the agents' expected utilities $J_0[I]$ and $J_1[I]$ are concave in I therefore, as was indicated in Proposition 1, optimization of the weighted sum (24) at each fixed $\delta > 0$ yields the set of all Pareto-optimal policies. The concavity is not the case if $R(y)$ is nonlinear, and the procedure gives then only a subset of all Pareto-optimal policies. If the mean value principle is used, i.e., $P = (1 + \alpha)E I(X)$, we have $R'(y) = 1 + \alpha > 1$, hence Proposition 3 is applicable and divides all the Pareto-optimal policies $\{\hat{I}_\delta(x)\}_{\delta > 0}$ into upper limit and deductible policies in terms of the weight δ . Thus, the premium size function $R(y) = (1 + \alpha)y$ turns out to be the "most convenient" for the analysis.

As was shown in Raviv (1979, section III) and Aase (2002, p. 108), any Pareto-optimal policy in the classical model (with zero administrative costs), where a premium size P is chosen jointly with $I(x)$, is a coinsurance policy with neither a deductible nor an upper limit. However, it is not the case in our model. Proposition 3 provides some explanation for deductible applications in practice even if administrative costs are negligible. It states that any Pareto-optimal policy corresponding to a sufficiently small relative insurer's weight δ ($\delta < \delta^*$) has a deductible $\hat{x} > 0$. A prevalence of deductible treaties over upper limit treaties in practice can thus be explained by the fact that an insurer is often not "very influential" (his δ is less than the critical weight δ^*) in making the contract with a potential policyholder.

Example 4: As in Example 1, we consider the case of quadratic utility functions $u_i(x) = -\frac{1}{2}\beta_i x^2 + x$, $i = 0, 1$. (Note that the domains are the intervals $(-\infty, 1/\beta_i)$ so it

is natural to assume that the initial capitals $w_i < 1/\beta_i, i = 0, 1$. Let the premium size function $R(y)$ be linear, $R(y) = (1 + \alpha)y$, with a given loading coefficient $\alpha > 0$.

As $u_i''(x) \equiv -\beta_i$, the right-hand side of (26) becomes the constant $\theta = \beta_1/(\delta\beta_0 + \beta_1)$ and the Pareto-optimal $\hat{I}(x)$ is either of the form $I_1(\hat{x}, x) = x - (1 - \theta)(x - \hat{x})_+$ or of the form $I_2(\hat{x}, x) = \theta(x - \hat{x})_+$. To determine the form of $\hat{I}(x)$ under a given δ , observe that the policy $I_0(\delta, x) = \theta x$ and, therefore, Equation (33) in Proposition 3 takes here the form

$$\delta(\beta_0 w_0 - 1) + 1 - \beta_1 w_1 + (1 + \alpha)\beta_1 E X = 0.$$

Its unique solution, the "critical weight," is

$$\delta^* = \frac{1 - \beta_1(w_1 - (1 + \alpha)EX)}{1 - \beta_0 w_0} > 0. \quad (34)$$

The last inequality holds as $w_i < 1/\beta_i, i = 0, 1$. According to Proposition 3, if the given insurer's relative weight $\delta < (=)\delta^*$ then the Pareto-optimal $\hat{I}(x)$ is a policy $I_2(\hat{x}, x)$ with a deductible $\hat{x} > (=)0$. Otherwise, $\hat{I}(x)$ is a policy $I_1(\hat{x}, x)$ with an upper limit $\hat{x} > 0$. To find, in every case, the value of $\hat{x}$, we make use of optimality condition (27). Observing that in the case of an upper limit policy

$$E I_1(\hat{x}, X) = \frac{\beta_1}{\delta\beta_0 + \beta_1} \int_{\hat{x}}^T \bar{F}(x) dx, \text{ where } \bar{F}(x) \stackrel{\text{def}}{=} 1 - F(x),$$

Equation (27) is equivalent to

$$\hat{x} - (1 - \alpha^2) \int_0^{\hat{x}} \bar{F}(x) dx + a = 0, \quad (35)$$

$$\text{where } a = \alpha \left\{ w_0 - \beta_0^{-1} + \frac{1 - \beta_1(w_1 - (1 + \alpha)EX)}{\delta\beta_0} \right\}.$$

Similarly, observing that in the deductible case

$$E I_2(\hat{x}, X) = EX - \frac{\delta\beta_0}{\delta\beta_0 + \beta_1} \int_{\hat{x}}^T \bar{F}(x) dx,$$

(27) is reduced to

$$-\hat{x} + (\alpha^2 - 1) \int_{\hat{x}}^T \bar{F}(x) dx + b = 0, \quad (36)$$

$$\text{where } b = (1 + \alpha)EX - \alpha \left\{ w_1 - \beta_1^{-1} + \frac{\delta(1 - \beta_0 w_0)}{\beta_1} \right\}.$$

It is noticeable that if $\alpha = 1$ —the loading equals the risk premium—then, as follows from (34)–(36), the Pareto-optimal policy $\hat{I}(\hat{x}, x)$ depends on the loss distribution $F(x)$ through the mean EX only, regardless of its shape.

For instance, let $\beta_0 = \beta_1 = 1/2$, $w_0 = 1$, $w_1 = 0$, $\alpha = 1$, the loss X have the mean $1/2$ (say, X is uniformly distributed on $[0, 1]$), and the weight $\delta > 0$ be a parameter. Then from (34) we get the “critical weight” $\delta^* = 3$. Let $\delta \leq 3$, then the Pareto-optimal $\hat{I}(x)$ is a deductible policy $(\delta + 1)^{-1}(x - \hat{x})_+$. Equation (36) gives $3 - \delta - \hat{x} = 0$ and, noting that the left-hand side is positive on $[0, 1]$ for $\delta < 2$, we obtain $\hat{x} = \min\{3 - \delta, 1\}$ (see Remark 3). If $\delta > 3$ then $\hat{I}(x)$ is an upper limit policy $x - \delta/(\delta + 1)(x - \hat{x})_+$ and from (35) we obtain $\delta(\hat{x} - 1) + 3 = 0$, therefore $\hat{x} = 1 - 3/\delta$. To sum up, if δ increases from 3 then the upper limit $\hat{x}$ increases from 0 and becomes 1 in the limit as $\delta \rightarrow \infty$; when δ decreases from 3 to 2, the deductible $\hat{x}$ grows from 0 to 1 and any further decrease in δ does not change $\hat{x} = 1$. This extreme deductible value means that the insurer does not take any risk, i.e., no contract is made under $\delta \leq 2$.

Consider the same problem with the exception that now $\delta = 3$ is fixed and the loading coefficient is increased to the value $\alpha = 1.5$. Equation (34) shows that the critical weight also increases to $\delta^* = 3.25 > \delta = 3$ and, hence, the Pareto-optimal $\hat{I}(x)$ is a deductible policy. Since for the uniform distribution $F(x)$ we have $\int_{\hat{x}}^1 \bar{F}(x) dx = (1 - \hat{x})^2/2$, (36) converts into not a linear but a quadratic equation $0.625(1 - \hat{x})^2 - \hat{x} - 0.25 = 0$, whence the optimal deductible is $\hat{x} = 0.175$.

APPENDIX

Proof of Proposition 2: Our proof partially follows the reasonings used in Raviv (1979, Lemma 1). Since the right-hand side of Equation (6) is decreasing in δ , it can be seen that its solution $\hat{I}(x)$ on $[\hat{x}, T]$ is also decreasing in δ for $x > \hat{x}$ under any of the initial conditions, $\hat{I}(\hat{x}) = \hat{x}$ or $\hat{I}(\hat{x}) = 0$. Indeed,

$$\frac{\partial \hat{I}(x)}{\partial \delta} = \int_{\hat{x}}^x \frac{\partial}{\partial \delta} \hat{I}'(t) dt = \int_{\hat{x}}^x \frac{\partial \hat{I}'(t)}{\partial \hat{I}} \frac{\partial \hat{I}(t)}{\partial \delta} + \frac{\partial \hat{I}'(t)}{\partial \delta} dt.$$

After differentiating $\frac{\partial \hat{I}(x)}{\partial \delta}$ with respect to x , we get a linear nonhomogeneous differential equation of the first order with initial condition $\frac{\partial \hat{I}(\hat{x})}{\partial \delta} = 0$. Denote $W(x) = \int_{\hat{x}}^x \frac{\partial \hat{I}'(t)}{\partial \hat{I}} dt$, then the solution to the differential equation has the form

$$\frac{\partial \hat{I}(x)}{\partial \delta} = \exp[W(x)] \int_{\hat{x}}^x \frac{\partial \hat{I}'(t)}{\partial \delta} \exp[-W(t)] dt < 0$$

since from (6) it follows that the integrand is negative.

Show that the above-examined policy $\hat{I}(x)$ with a deductible (with an upper limit) is decreasing (increasing) in $\hat{x}$ for each $x > \hat{x}$. Similar to the reasonings above, we have in the first case

$$\hat{I}(x) = \int_{\hat{x}}^x \hat{I}'(t) dt, \text{ hence } \frac{\partial \hat{I}(x)}{\partial \hat{x}} = -\hat{I}'(\hat{x}) + \int_{\hat{x}}^x \frac{\partial \hat{I}'(t)}{\partial \hat{I}} \frac{\partial \hat{I}(t)}{\partial \hat{x}} dt.$$

Solving the latter equation gives

$$\frac{\partial \hat{I}(x)}{\partial \hat{x}} = -\hat{I}'(\hat{x}) \exp \left\{ \int_{\hat{x}}^x \frac{\partial \hat{I}(t)}{\partial \hat{I}} dt \right\} < 0.$$

The case of the policy with an upper limit is treated analogously.

Suppose first that the given $\delta > \delta^*$. Then, by the proved monotonicity in δ , $I_0(\delta, x) < I_0(\delta^*, x)$ and $E I_0(\delta, X) < E I_0(\delta^*, X) = M$ so that the policy $I_0(\delta, x)$ should be increased when changed to the Pareto-optimal $\hat{I}(x)$. By Theorem 1, $\hat{I}(x)$ must be of form either (3) or (4) and, as was proved, only policy (3) of them increases in $\hat{x}$. This shows that the Pareto-optimal $\hat{I}(x)$ is a policy with an upper limit. If $\delta < \delta^*$ then $E I_0(\delta, X) > M$. Taking into account that policy (4) is decreasing in $\hat{x}$, we come to the conclusion that $\hat{I}(x)$ is a policy with a deductible.

Prove the last statement of Proposition 2. Denote by $I_1(\hat{x}, x)$ the Pareto-optimal policy with an upper limit $\hat{x} > 0$. We have proved above that $I_1(\hat{x}, x)$ is a continuous increasing function of $\hat{x}$ on the interval $[0, T]$. Hence, the expectation $E I_1(\hat{x}, X)$ also continuously increases from $E I_0(\delta, X) < M$ to $E X > M$ when $\hat{x}$ changes from 0 to T . This, clearly, implies that the optimal upper limit $\hat{x}$ is a unique solution to the equation $E I_1(\hat{x}, X) = M$.

The case of the Pareto-optimal policy $I_2(\hat{x}, x)$ with a deductible $\hat{x}$ is considered in the same way. Q.E.D.

Proof of Proposition 3: Since δ^* is a root of (33) then, as was shown in Theorem 2, $\hat{U}(x) - \hat{C} \equiv 0$ on $[0, T]$ under $\hat{I}(x) = I_0(\delta^*, x)$ and, hence, $I_0(\delta^*, x)$ is optimal in Equation (24) where δ is taken equal to δ^* .

Let, first, the prescribed weight δ be such that $\delta < \delta^*$. Suppose the contrary to Proposition 3's statement, that is, the optimal $\hat{I}(x)$ is a policy with an upper limit $\hat{x} > 0$. As was shown in Theorem 2's proof, this means that $\hat{U}(0) - \hat{C} < 0$. If we continuously decrease δ up to zero, the corresponding optimal policy $\hat{I}(\delta, x)$ continuously changes to $\hat{I}(0, x)$, a solution to (24) with $\delta = 0$. By the definition (see Equation (29)), the expression for $\hat{U}(0) - \hat{C}$ with $\delta = 0$ takes the form

$$R'(E\hat{I}(0, X))Eu'_1(w_1 - \hat{P} - X + \hat{I}(0, X)) - u'_1(w_1 - \hat{P}) \quad (\text{A1})$$

where $\hat{P} = R(E\hat{I}(0, X))$. Observing that $R'(y) > 1$, $-x + \hat{I}(0, x)$ is a nonincreasing function, and $u'_1(x) < 0$, we know that expression (A1) is positive. Therefore there exists a point δ' such that $\hat{U}(0) - \hat{C} = 0$ for $\delta = \delta'$ and, hence, the optimal $\hat{I}(\delta', x)$ is a policy with neither upper limit nor deductible. As was proved above, it implies that $\delta' \in (0, \delta^*)$ is a solution to (33), but this contradicts the assumption that δ^* is the only solution to (33). The obtained contradiction proves that the policy $\hat{I}(x)$ we have started from is a policy with a deductible.

Let $\delta > \delta^*$. Following the reasonings above, we obtain in this situation that $\hat{I}(x)$ is a policy with an upper limit $\hat{x} > 0$. Q.E.D.

REFERENCES

- Aase, K. K., 2002, Perspectives of Risk Sharing, *Scandinavian Actuarial Journal*, 2: 73-128.
- Arrow, K., 1971, *Essays in the Theory of Risk Bearing* (Chicago: Markham).
- Baton, B., and J. Lemaire, 1981, The Core of a Reinsurance Market, *ASTIN Bulletin*, 12: 57-71.
- Blazenko, G., 1985, Optimal Indemnity Contracts, *Insurance: Mathematics and Economics*, 4: 267-278.
- Borch, K., 1960, Reciprocal Reinsurance Treaties, *ASTIN Bulletin*, 1: 171-191.
- Borch, K., 1990, *Economics of Insurance*, edited by K. K. Aase, and A. Sandmo (Amsterdam: North-Holland).
- Gerber, H. U., 1978, Pareto-Optimal Risk Exchanges and Related Decision Problems, *ASTIN Bulletin*, 10: 25-33.
- Golubin, A. Y., 2002, Franchise Optimization in the Static Insurance Model, *Journal of Mathematical Sciences*, 112: 4126-4140.
- Holmstrom, B., 1979, Moral Hazard and Observability, *Bell Journal of Economics*, 10: 74-91.
- Kalai, E., and M. Smorodinsky, 1975, Other Solutions to the Nash Bargaining Problem, *Econometrica*, 43: 513-518.
- Karlin, S., and W. Studden, 1966, *Tchebycheff Systems: With Applications in Analysis and Statistics* (New York: Interscience Publishers).
- Lehmann, E. L., 1959, *Testing Statistical Hypotheses* (New York: Wiley and Sons).
- Lemaire, J., 1979, A Non Symmetrical Value For Games Without Transferable Utilities. Application to Reinsurance, *ASTIN Bulletin*, 10: 195-214.
- Lemaire, J., 1990, Cooperative Game Theory and Its Insurance Applications, *ASTIN Bulletin*, 21: 17-40.
- Murray, M. L., 1971, A Deductible Selection Model: Development and Applications, *Journal of Risk and Insurance*, 38: 423-436.
- Nash, J., 1950, The Bargaining Problem, *Econometrica*, 18: 155-162.
- Pardalos, P. M., Y. Siskos, and C. Zopounidis, 1995, *Advances in Multicriteria Analysis* (Dordrecht: Kluwer Academic Publishers).
- Raviv, A., 1979, The Design of an Optimal Insurance Policy, *American Economic Review*, 69: 84-96.
- Schlesinger, H., 1981, The Optimal Level of Deductibility in Insurance Contracts, *Journal of Risk and Insurance*, 48: 465-481.
- Wyler, E., 1990, Pareto Optimal Risk Exchanges and a System of Differential Equations: A Duality Theorem, *ASTIN Bulletin*, 20: 23-32.

Copyright of Journal of Risk & Insurance is the property of Blackwell Publishing Limited and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.