

REINSURANCE ARRANGEMENTS MAXIMIZING INSURER'S SURVIVAL PROBABILITY

Lesław Gajek
Dariusz Zagrodny

ABSTRACT

The article concerns the problem of purchasing a reinsurance policy that maximizes the survival probability of the insurer. Explicit forms of the contracts optimal for the insurer are derived which are stop loss or truncated stop loss depending on the initial surplus, a quota to be spend on reinsurance and pricing rules of both the insurer and the reinsurer.

INTRODUCTION

A standard result concerning a purchase of an optimal insurance policy is given by Arrow (1963), who examined in detail the economic significance of medical industry in the United States. In that paper, the optimal policy was understood as the one that maximized the expected utility of the buyer who was risk averse (i.e., with concave utility function). In the present article, we reconsider the problem of choosing an optimal policy in another important special case: when the policy buyer is an insurance company. Neglecting the tax regulations, a clear aim for such risk protection is decreasing the probability of ruin or, equivalently, increasing the probability of survival. However, this appears to be quite difficult (see Kaas et al., 2001, p. 92). Indeed, the probability of survival is the expectation of an indicator function, which is neither convex nor concave. To be more precise, let X , a random variable, denote the total claim amount during the considered period of time. Let π_X denote the collected premiums from the insureds during the same time and let $R(X)$ be the part of total claims that is covered by the reinsurer (a more detailed discussion of measurability assumptions on R can be found in "Optimal Partial Protection Against Ruin"). Following Arrow (1963), we shall assume throughout the article that

$$0 \leq R(X) \leq X$$

almost everywhere and that the insurer has to pay for the reinsurance protection the premium $\pi_R \equiv \pi_{re}(ER(X))$, where $\pi_{re}(\cdot)$ is an increasing function. Throughout the

Lesław Gajek works at the Institute of Mathematics, Technical University of Łódź, Wólczńska 215, 93-005 Łódź, Poland, e-mail: gal@p.lodz.pl. Dariusz Zagrodny works at the Faculty of Mathematics, Cardinal Stefan Wyszyński University, Dewajtis 5, 01-815 Warsaw, Poland. This work was supported by KBN grant no. 2 H02B O11 24.

article we assume that $EX < \infty$ and will denote the initial surplus of the insurer by V . The eventual insurer's surplus, $Y(X)$, resulting from the reinsurance contract R , can be expressed as follows:

$$Y(X) = V + \pi_X - \pi_R + R(X) - X.$$

The ruin happens whenever $Y(X) < 0$ so the probability of survival, $\varphi(R)$, can be expressed in the following way:

$$\varphi(R) = E\mathbb{I}_{[0,\infty)}(V + \pi_X - \pi_R + R(X) - X).$$

Note that the indicator $\mathbb{I}_{[0,\infty)}(\cdot)$ is not a concave function (while it is nondecreasing) and the initial surplus $V + \pi_X - \pi_R$ depends on R thus the problem considered here cannot be reduced to the one considered in Arrow (1963). According to proposition 1 of Arrow (1963), the optimal policy for the insurer would be the so-called *stop loss* reinsurance:

$$R(t) = (t - M)^+,$$

where M is a deductible and a^+ means $\max(0, a)$ for any real a . But such a contract provides either the full protection against the ruin, when

$$V + \pi_X - \pi_R \geq M, \tag{1}$$

or it does not decrease the ruin probability at all, when the opposite inequality holds. In the latter case, the stop loss protection is useless because the insurer is ruined before it starts working. This analysis leads to the following two important questions:

- (a) If we decide to purchase a full protection against ruin, is the stop loss policy the least expensive one?
- (b) If we look for a partial protection only, what is the best reinsurance contract at a price not exceeding a given quota P that we have decided to spend on it?

In "Optimal Full Protection Against Ruin," we first show that indeed the answer to (a) is affirmative. In the next section, we formulate and prove a proposition that concerns the explicit form of the optimal reinsurance contract in case (b). If randomized decisions are not allowed, it turns out that the optimal contract is of the following form

$$R^*(t) = \mathbb{I}_K(t)(t - V - \pi_X + P)^+ \tag{2}$$

with some set $K \subseteq [0, \infty)$, which depends on both P and the initial assets $V + \pi_X - P$ (throughout the article we assume that the initial assets are nonnegative). An interesting observation is that randomized decisions can produce lower ruin probability than nonrandomized ones, under a given upper bound P on the price of the reinsurance policy (see Example 4). However, the improvement due to randomization diminishes with decreasing the probability mass of every atom. Another way

to eliminate randomized decisions is an adjustment of the upper limit price P (see Example 5).

Randomization cannot be helpful if the total claim distribution is absolutely continuous. Then the optimal contract R^* takes the form

$$R^*(t) = \begin{cases} 0 & \text{when } 0 \leq t \leq V + \pi_X - P, \\ t - V - \pi_X + P & \text{when } V + \pi_X - P < t < c^*, \\ 0 & \text{when } t \geq c^*, \end{cases}$$

where c^* is such a constant that $\pi_{\text{re}}(ER^*) = P$. It should be stressed, however, that the optimality of R^* is subject to some regularity Condition (14), which cannot be dropped (see Example 3 for strange consequences of violating it).

As a byproduct, we reveal in “Optimal Partial Protection Against Ruin” a beautiful link between purchasing an optimal reinsurance policy and testing statistical hypothesis (TSH) that concurs with the Arrow’s analysis of the theory of choice in risk-taking situations (see Arrow, 1951).

To summarize, the contract that maximizes survival probability is not, in general, of the same form as in the case of maximizing expected concave utility. For a large value of the initial surplus it may coincide with a stop loss reinsurance, depending on the quota P , but for small ones it is rather a truncated stop loss (2) (or its randomized version). A more detailed analysis can be found in “Optimal Partial Protection Against Ruin” and “Numerical Examples.” The last section concludes the article.

Reinsurance arrangements that minimize the variance of the insurer’s portfolio under other pricing rules can be found in Gajek and Zagrodny (2000), Mazur (2000), and Kaluszka (2001). In Gajek and Zagrodny (2004), a general risk measure of the insurer is minimized after a purchase of reinsurance arrangement priced according to the standard deviation rule. Amsler (1993) proposed a reinsurance arrangement under which the reinsurer provides a full protection against ruin during a contractual period. The reinsurer’s share is adjusted dynamically to cover the actual deficit; however, this treaty was not proved to be optimal in our sense. Other results on optimal reinsurance are reviewed, e.g., in Kass et al. (2001), Daykin, Pentikäinen, and Pesonen (1993), Pesonen (1984), and Verlaak and Beirlant (2003). In the latter paper, one can find also some comments on a common reinsurance practice.

OPTIMAL FULL PROTECTION AGAINST RUIN

It is easy to notice that if $\Pr(X \leq V + \pi_X) = 1$, we have no possibility of ruin without any reinsurance protection. So the only interesting case is when

$$\Pr(X > V + \pi_X) > 0, \quad (3)$$

and we shall consider this case throughout the article.

Next, let us observe that if

$$\pi_{\text{re}}(E(X + P - \pi_X - V)^+) \leq P, \quad (4)$$

we have enough money to purchase the stop loss policy

$$(X - M_P)^+,$$

with the deductible $M_P = V + \pi_X - P$, which provides the full protection against ruin. Indeed, after having bought this policy the probability of ruin becomes zero. Whether (4) holds depends very much on V in the way that for sufficiently large V (4) can always be satisfied except for the case when $\pi_{re}(0) > P$, if π_{re} is continuous at 0. However, it is quite natural to assume that $\pi_{re}(0) = 0$, which means that the contract without any transfer of risk would cost nothing. It is also quite natural that

$$\pi_{re}(EX) < V + \pi_X. \quad (5)$$

Condition (5) may be interpreted that the price of transferring of the whole portfolio to the reinsurer does not exceed the collected premium amount plus the initial surplus. If (5) is not satisfied it means that the reinsurer uses a completely different pricing rule than the insurer does because otherwise $\pi_X \cong \pi_{re}(EX)$ and (5) automatically holds even for relatively small values of V . Another explanation for violating (5) is that the insurer sells the policies below their true value and accepts an extraordinary risk due to this "dumping."

Let us notice that each of the Conditions (4) and (5) implies that for some $M \geq 0$ the following inequality holds:

$$\pi_{re}(E(X - M)^+) \leq V + \pi_X - M. \quad (6)$$

Proposition 1: Assume that M^* is the largest value of M for which (6) holds. Then the stop loss

$$R^*(x) = (x - M^*)^+$$

is the least-expensive reinsurance arrangement with the probability of ruin equal to 0.

Proof: Since $(x - M)^+$ is decreasing in M , the same is true for $E(X - M)^+$, which implies that the price of stop loss decreases when M increases because π_{re} is increasing. If there are several such M that satisfy (6), the largest one, M^* , corresponds to the lowest price of the stop loss policy. We shall show that for every reinsurance rule R such that the inequalities $0 \leq R(X) \leq X$ and

$$X - R(X) + \pi_{re}(ER) - V - \pi_X \leq 0, \quad (7)$$

hold with probability 1, the price of R cannot be smaller than the price of $R^*(x) = (x - M^*)^+$. So let R be such that (7) holds almost surely and define

$$M_R = \operatorname{ess\,sup}_{x \geq 0} (x - R(x)), \quad (8)$$

where the essential supremum is with respect to the claim distribution. Clearly $M_R < \infty$ because otherwise (7) almost surely could not hold. It is also enough to consider

the case $M_R > 0$, because $\pi_{\text{re}}(E(X - M^*)^+) \leq \pi_{\text{re}}(EX)$ and otherwise the reinsurance R coinciding almost everywhere with X could not be less expensive than the stop loss with retention M^* . By (7) and the definition of M_R ,

$$M_R + \pi_{\text{re}}(ER) - V - \pi_X \leq 0. \quad (9)$$

Let us denote $R_M(x) = (x - M_R)^+$ and observe that $R_M(x) \leq R(x)$ almost everywhere. Hence $ER_M \leq ER$ and $\pi_{\text{re}}(ER_M) \leq \pi_{\text{re}}(ER)$. So the stop loss R_M is less expensive than R and

$$\Pr(X - R_M(X) + \pi_{\text{re}}(ER_M) - V - \pi_X \leq 0) = 1$$

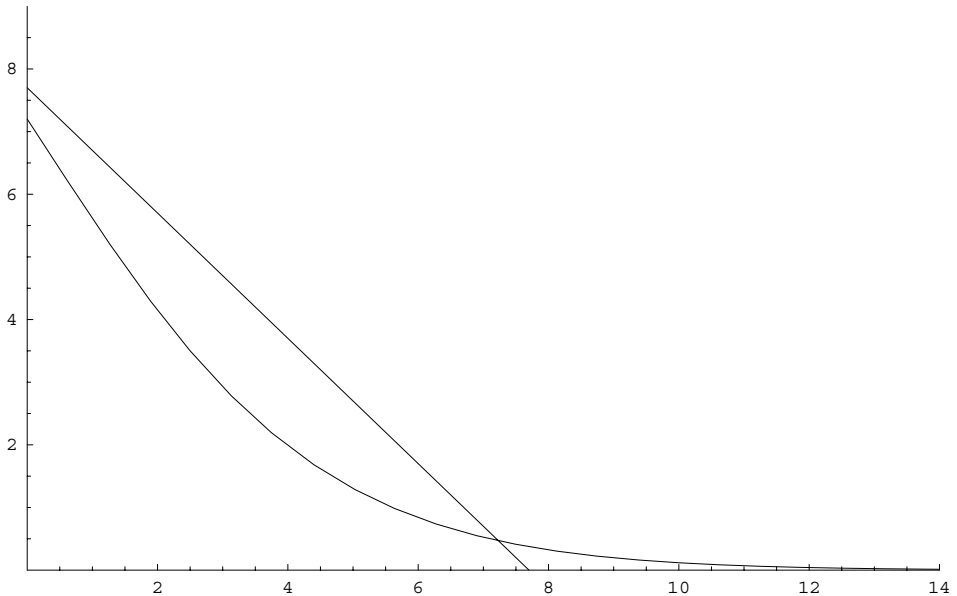
by (9), (8), and the definition of R_M . But R^* is the least-expensive stop loss satisfying (7) so the price of R cannot be smaller than that of R^* . Q.E.D.

In order to investigate, if M^* exists, let us assume that π_{re} is a continuous function that takes positive values on $(0, \infty)$. If (5) holds, then the graph of $\pi_{\text{re}}(E(X - M)^+)$, as a function of M , typically looks like a curved line in Figure 1.

Observe that both π_{re} and the linear function on the right side of (6) (with graphs demonstrated in Figure 1) are continuous in M (continuity of $E(X - M)^+$ in M is true

FIGURE 1

Graphs of $\pi_{\text{re}}(E(X - M)^+)$ (Curved Line) and the Right Side of (6) (Straight Line), as Functions of M , for Gamma Distributed Total Claims; for Details See Example 1



also for discrete distributions of X). Hence there is a deductible $M^* < V + \pi_X$ such that

$$M \leq V + \pi_X - \pi_{\text{re}}(E(X - M)^+) \quad \text{for } M \leq M^* \quad (10)$$

and

$$M > V + \pi_X - \pi_{\text{re}}(E(X - M)^+) \quad \text{for } M > M^*, \quad (11)$$

provided that for $M = V + \pi_X$ the left side of (6) is positive. In view of (3), this is true if $\pi_{\text{re}}(t) > 0$ for every positive t .

From (11), we realize that the stop loss reinsurance with $M > M^*$ is useless because it never starts working before the insurer is ruined. On the other hand, Inequality (10) means that stop loss with any $M \leq M^*$ gives a full protection against ruin because ruin cannot occur. Considering $\pi_{\text{re}}(E(X - M)^+)$ is decreasing in M , we arrive at the conclusion that the least expensive stop loss contract that provides the full protection against ruin is that one with the deductible M^* defined as a unique solution of the equation

$$M^* = V + \pi_X - \pi_{\text{re}}(E(X - M^*)^+). \quad (12)$$

That is, at least one such M^* exists because of continuity of $\pi_{\text{re}}(\cdot)$ and the facts that $\pi_{\text{re}}(t) > 0$ on $(0, \infty)$ and that (6) is satisfied for at least one $M > 0$. The latter condition is fulfilled, e.g., when one of the Inequalities (4) and (5) holds. Throughout the article we shall assume that $\pi_{\text{re}}(EX) < V + \pi_X$.

It may also happen that there are several $M > M_p$ for which (6) holds, i.e.,

$$M \leq V + \pi_X - \pi_{\text{re}}(E(X - M)^+).$$

We will discuss this case in detail in Example 2.

OPTIMAL PARTIAL PROTECTION AGAINST RUIN

Since π_{re} takes positive values on $(0, \infty)$ and $\Pr(X > V + \pi_X) > 0$, Inequality (4) cannot hold for all $P > 0$. There always exists a sufficiently small quota $P > 0$ such that (4) is no longer satisfied. So consider now the case when

$$\pi_{\text{re}}(E(X - V - \pi_X + P)^+) > P \quad (13)$$

for all sufficiently small P . This means that even spending the whole quota P on protection against ruin, we are not able to purchase a stop loss policy that decreases the probability of ruin to zero.

To prove that the truncated stop loss R^* defined by (2) is optimal, we shall use a completely different technique than Arrow did. Surprisingly enough, it turns out that the decision problem considered here is closely related to finding the most powerful test in testing a simple statistical hypothesis against a simple alternative. Therefore

we can apply the fundamental Neyman–Pearson lemma (see, e.g., Lehmann, 1985) to construct the optimal reinsurance contract. Thus Proposition 2 and Corollary 1 reveal the link between the above two possible approaches to choose an optimal decision and thus completes the analysis of Arrow's (1951) theory of choice in risk-taking situations in this case.

We shall start with formulating a conjugate problem of TSH. Let $v \equiv V + \pi_X - P$. Consider two probability measures on $(v, +\infty)$, Q_0 , and Q_1 , such that

$$dQ_0(y) = \frac{(y-v)dF(y)}{(1-F(v))e(v)} \quad \text{and} \quad dQ_1(y) = \frac{dF(y)}{1-F(v)},$$

where F denotes the total claim distribution function and $e(v)$ is the mean excess loss for a deductible of v , i.e., the expected loss X in excess of the initial assets v conditioned on $X > v$ (see Klugman, Panjer, and Willmot, 1998).

Let Y be a random variable taking values in $(v, +\infty)$. Consider the following conjugate problem of TSH:

H_0 : Y has a probability distribution Q_0

versus

H_1 : Y has a probability distribution Q_1 .

Let K_α denote the most powerful critical set, i.e., $Q_0(K_\alpha) \leq \alpha$ and $Q_1(K_\alpha)$ take the largest possible value over the class of all the sets with the Q_0 measure not exceeding α . In Proposition 2, we restrict our attention to nonrandomized tests and reinsurance rules only.

Proposition 2: Suppose that $\pi_{\text{re}}(\cdot)$ is strictly increasing and

$$\pi_{\text{re}}(t + (1 - F(V + \pi_X - P))(P - \pi_{\text{re}}(t))) \leq P \quad (14)$$

for all $t \in [0, \pi_{\text{re}}^{-1}(P))$. Then

$$R^*(t) = \mathbb{I}_K(t)(t - V - \pi_X + P)$$

is optimal, if $K = K_\alpha$ is the most powerful critical set in the conjugate TSH problem with

$$\alpha = \frac{\pi_{\text{re}}^{-1}(P)}{[1 - F(V + \pi_X - P)]e(V + \pi_X - P)}. \quad (15)$$

Proof: First of all, α given by (15) is smaller than 1, by (13). Next, observe that

$$0 \leq R^*(x) \leq x$$

and

$$Q_0(K_\alpha) \leq \frac{\pi_{\text{re}}^{-1}(P)}{[1 - F(V + \pi_X - P)]e(V + \pi_X - P)},$$

which means that

$$\int_{K_\alpha} (y - V - \pi_X + P) dF(y) = ER^* \leq \pi_{\text{re}}^{-1}(P).$$

Since $\pi_{\text{re}}(\cdot)$ is increasing,

$$\pi_{\text{re}}(ER^*) \leq P$$

and R^* is feasible. Now let us denote the ruin probability $1 - \varphi$ by ψ and observe that

$$\begin{aligned} \psi(R^*) &= \int_{[0, V + \pi_X - P]} \mathbb{I}_{(0, \infty)}(x - [V + \pi_X - \pi_{R^*}]) dF(x) \\ &\quad + \int_{K_\alpha} \mathbb{I}_{(0, \infty)}(V + \pi_X - P - [V + \pi_X - \pi_{R^*}]) dF(x) \\ &\quad + \int_{(V + \pi_X - P, \infty) \setminus K_\alpha} \mathbb{I}_{(0, \infty)}(x - [V + \pi_X - \pi_{R^*}]) dF(x) \\ &\leq Q_1((v, \infty) \setminus K_\alpha)[1 - F(v)]. \end{aligned} \tag{16}$$

If for some feasible R we have

$$\psi(R) < \psi(R^*), \tag{17}$$

then define the set K as follows:

$$K = \{x \in (v, \infty) \mid x - R(x) - (V + \pi_X - \pi_R) \leq 0\},$$

and observe that

$$\begin{aligned} &Q_1((v, \infty) \setminus K)[1 - F(v)] \\ &= \int_{(v, \infty) \setminus K} \mathbb{I}_{(0, \infty)}(x - R(x) - [V + \pi_X - \pi_R]) dF(x) \\ &\leq \int_{(0, \infty)} \mathbb{I}_{(0, \infty)}(x - R(x) - [V + \pi_X - \pi_R]) dF(x) = \psi(R) < \psi(R^*) \quad \text{by (17)} \\ &\leq Q_1((v, \infty) \setminus K_\alpha)[1 - F(v)] \quad \text{by (16)}. \end{aligned}$$

If $ER < \pi_{\text{re}}^{-1}(P)$, (14) implies the inequalities

$$\begin{aligned} &\pi_{\text{re}}(ER + E\mathbb{I}_K(X)(P - \pi_R)) \\ &\leq \pi_{\text{re}}(ER + (1 - F(V + \pi_X - P))(P - \pi_R)) \leq P. \end{aligned}$$

By definition, for every $x \in K$

$$R(x) + P - \pi_R \geq x - (V + \pi_X - \pi_R) + P - \pi_R = x - v.$$

This, however, would contradict the fact that K_α is a most powerful critical set at the significance level α because

$$Q_0(K) = \int_K \frac{(y - v) dF(y)}{[1 - F(v)]e(v)} \leq \frac{ER + E\mathbb{I}_K(x)(P - \pi_{re})}{[1 - F(v)]e(v)} \leq \frac{\pi_{re}^{-1}(P)}{[1 - F(v)]e(v)} = \alpha. \quad (18)$$

If $ER = \pi_{re}^{-1}(P)$, Inequality (18) follows directly from the inequality $R(x) \geq x - v$, which holds on K . So in both cases we arrive at a contradiction, which completes the proof. Q.E.D.

At first glance, Assumption (14) seems to be difficult to interpret. However, in case when $\pi_{re}(\cdot)$ is linear with a safety loading θ_{re} , Assumption (14) is satisfied if

$$\frac{\theta_{re}}{1 + \theta_{re}} \leq F(V + \pi_X - P). \quad (19)$$

The left side of (19) is equal to the infinite horizon survival probability of the reinsurer with no initial reserve in the compound Poisson model (see Klugman, Panjer, Willmot, 1998, p. 536). Thus, (14) admits an easy interpretation in terms of the survival probabilities of the insurer and the reinsurer. Consider now a general case when π_{re} can be nonlinear and suppose that π_{re} is Lipschitzian with a Lipschitz constant L . Then, by the Tchebyshev inequality,

$$\begin{aligned} & \pi_{re}(t + (1 - F(V + \pi_X - P))(P - \pi_{re}(t))) - \pi_{re}(t) \\ & \leq L(1 - F(V + \pi_X - P))(P - \pi_{re}(t)) \\ & \leq \frac{LEX}{V + \pi_X - P}(P - \pi_{re}(t)). \end{aligned}$$

So if $L \leq (V + \pi_X - P)/EX$, Assumption (14) is satisfied. In fact, it is enough to assume that the highest rate of increase of π_{re} is bounded by $(V + \pi_X - P)/EX$. Since $\theta_X \equiv (\pi_X - EX)/EX$ can be interpreted as a global safety loading coefficient of the insurer, and

$$\frac{d}{dt}\pi_{re}(t) - 1 \equiv \theta_{re}(t)$$

as a marginal safety loading of the reinsurer at t , it means that the following relationship between the safety loading coefficients holds:

$$\sup_{0 \leq t < \pi_{re}^{-1}(P)} \theta_{re}(t) \leq \theta_X + \frac{V - P}{EX}, \quad (20)$$

where typically V is larger than P .

From the fundamental Neyman–Pearson lemma (see Lehmann, 1985) for every significance level α , $1 > \alpha > 0$, the most powerful test ϕ exists. More precisely, there are constants k and γ such that

$$\phi(y) = \begin{cases} 1 & \text{when } dQ_1(y)/dF(y) > kdQ_0(y)/dF(y), \\ \gamma & \text{when } dQ_1(y)/dF(y) = kdQ_0(y)/dF(y), \\ 0 & \text{when } dQ_1(y)/dF(y) < kdQ_0(y)/dF(y), \end{cases} \quad (21)$$

where $0 \leq \gamma \leq 1$ is chosen in such a way that $\int \phi(y) dQ_0(y) = \alpha$. The randomization via $\gamma \in (0, 1)$ is not necessary, if F is absolutely continuous.

So, taking into account (21) and Proposition 2, we get the following corollary.

Corollary 1: *Suppose F is absolutely continuous with a density positive almost everywhere on $[0, \infty)$. Under the assumptions of Proposition 2, there is a constant $c^* > V + \pi_X - P$ such that $K_\alpha = (V + \pi_X - P, c^*)$. Therefore, the optimal reinsurance contract is of the form*

$$R^*(t) = \begin{cases} 0 & \text{when } 0 \leq t \leq V + \pi_X - P, \\ t - V - \pi_X + P & \text{when } V + \pi_X - P < t < c^*, \\ 0 & \text{when } t \geq c^*, \end{cases} \quad (22)$$

where the truncation constant c^* can be calculated from the equation $\pi_{\text{re}}(\text{ER}^*) = P$.

To prove Corollary 1 it is enough to observe that the constant k in (21) must be positive because $\alpha < 1$. Hence $K_\alpha = \{y > v : y < c^*\}$, where $c^* = v + k^{-1}e(v)$. The equation $\int \phi(x) dQ_0(x) = \alpha$, with α given by (15), reduces now to $\text{ER}^* = \pi_{\text{re}}^{-1}(P)$, which always has a solution because X is absolutely continuous and (13) holds.

So far we have assumed that R is a real Borel function defined on $[0, \infty)$. This assumption excludes randomized reinsurance strategies. However, if X has a discrete distribution function the most powerful test ϕ at a given significance level α may be randomized. It means that for some claims the hypothesis H_0 is rejected with a probability γ and accepted with a probability $1 - \gamma$ (see (21)), which admits dependence of R on additional argument. The randomization can enlarge the power of the test in case the exact value α cannot be achieved on the set K_α for any k . By the link revealed in Proposition 2 and Corollary 1, this corresponds to using randomized reinsurance strategies which rely on covering a limit claim $X = c^*$ in part $c^* - V - \pi_X + P$ by the reinsurer—with a probability γ , and by the insurer—with a probability $1 - \gamma$. However, the randomization has rather theoretical meaning: first of all, the total claim distribution can be usually approximated with a high accuracy by an absolutely continuous distribution; secondly, the sides of the contract can usually change properly the quota P in order to achieve the equality $\pi_{\text{re}}(\text{ER}^*) = P$ without any randomization (see the example given in “Numerical Examples”).

NUMERICAL EXAMPLES

Example 1: Assume that the monetary unit is 1 billion USD and the total claim is gamma distributed with the shape parameter 3 and the scale parameter 1.5 (for the definition and properties of gamma distribution, see Klugman, Panjer, and Willmot

(1998, p. 579)). The expected claim then equals to 4.5 billion USD. Assume that the total premium collected by the insurer $\pi_X = 7.2$ billion USD and that the initial surplus $V = 0.5$ billion USD. Assume further that the reinsurer uses the pricing rule

$$\pi_R = (1 + \theta_{re})ER$$

with the safety loading coefficient $\theta_{re} = 0.6$. In Figure 1, the price of a stop loss reinsurance $(X - M)^+$, as a function of M , is plotted along with the linear function $V + \pi_X - M$, which appears on the right side of (6). The retention M^* of the least expensive stop loss providing the full protection against ruin is equal to 7.23 billion USD. The price of this reinsurance is

$$\pi_{re}(E(X - M^*)^+) = 0.47 \text{ billion USD.}$$

If $P \geq 0.47$ billion USD, we can purchase this reinsurance policy that reduces the ruin probability of the insurer to 0 from the initial value 0.114 corresponding to the case with no reinsurance protection. If $P < 0.47$ billion USD, we are not able to purchase a reinsurance policy reducing the ruin probability to 0. Assume, for example, that $P = 0.2$ billion USD. Since $\theta_X = (\pi_X - EX)/EX = 0.6$ and $V > P$, Equation (20) obviously holds and so does (14). Therefore the contract optimal for the insurer is the truncated stop loss R^* given by (22) with the truncation constant c^* , which can be found from the equation

$$\pi_{re}(ER^*) = P.$$

Solving numerically this equation, we get $c^* = 10.80$ billion USD. After establishing this reinsurance arrangement, the ruin probability decreases from 0.114 to 0.026.

Now we shall show that Equation (12) may have several roots depending on the shape of pricing function π_{re} .

Example 2: Let the claim distribution be as in Example 1 and

$$\pi_{re}(t) = 0.875 \begin{cases} 2t & \text{when } 0 \leq t < 0.5, \\ 1 + 4(t - 0.5) & \text{when } 0.5 \leq t < 2, \\ 7 + 0.2(t - 2) & \text{when } t > 2. \end{cases} \quad (23)$$

The price of a stop loss reinsurance under this pricing rule is plotted in Figure 2. Observe that now there are three positive roots of Equation (12).

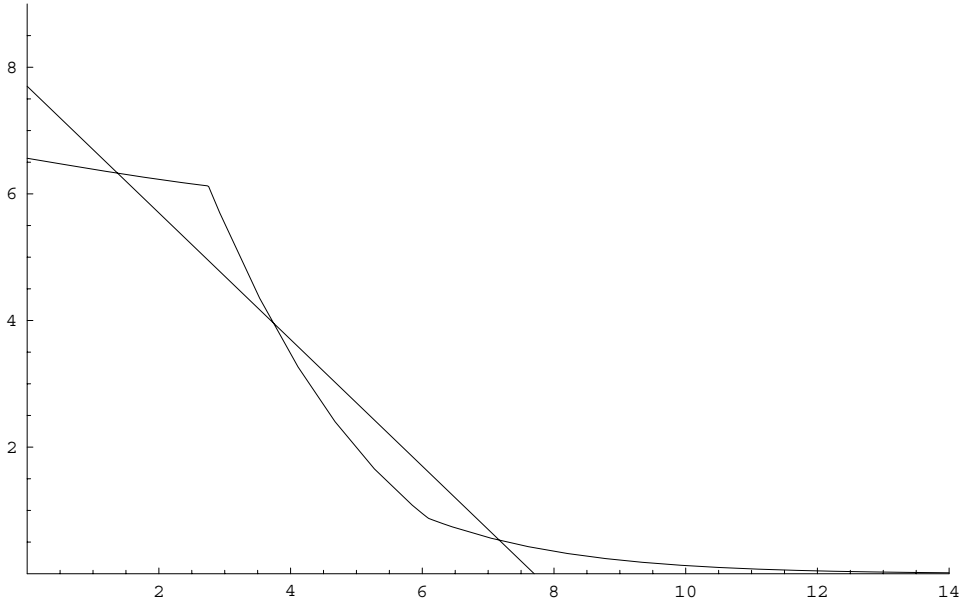
The next example shows strange consequences of dropping Assumption (14).

Example 3: Let the total claims have a distribution with the following Lebesgue density

$$f(x) = \begin{cases} 1 & \text{when } 1.5 \leq x \leq 2, \\ 0.25 & \text{when } 2 < x \leq 4, \\ 0 & \text{otherwise.} \end{cases}$$

FIGURE 2

Both Sides of (6) Plotted as Functions of M for Gamma Distributed Claims and π_{re} Defined by (23)



Further the pricing rule is defined by

$$\pi_{re}(t) = \begin{cases} 6t & \text{when } 0 \leq t \leq 0.25, \\ 1.5 + 0.33(t - 0.25) & \text{when } t > 0.25. \end{cases}$$

Let $P = 1.5$ billion USD, $V = 0.625$ billion USD, and $\pi_X = 2.375$ billion USD. It is easy to check that the left side of (14) is greater than P for all $t \in [0, \pi_{re}^{-1}(P))$, where $\pi_{re}^{-1}(P) = 0.25$, so (14) is not satisfied. Now observe that the truncated stop loss at price $P = 1.5$ billion USD has the truncation point $c_P = 2.62$ billion USD and the probability of ruin

$$\psi_P = 0.345.$$

But if we decide to spend only $P' = 1$ billion USD, we can purchase a truncated stop loss which has the truncation point $c_{P'} = 3.15$ billion USD and the probability of ruin

$$\psi_{P'} = 0.2125.$$

So the less expensive truncated stop loss contract would provide a higher protection against the ruin, if (14) is not assumed to hold. Further, if we do not purchase any protection, the probability of ruin will be 0.25, which is less than $\psi_P = 0.345$ that would cost 1.5 billion USD!

The next example concerns possible improvement of reinsurance policy by admitting randomization.

Example 4: Let X have a probability distribution

$$P(X = x) = \begin{cases} 0.5 & \text{for } x = 0.05 \text{ billion USD,} \\ 0.125 & \text{for } x = 0.1 \text{ billion USD,} \\ 0.175 & \text{for } x = 0.2 \text{ billion USD,} \\ 0.125 & \text{for } x = 0.3 \text{ billion USD,} \\ 0.075 & \text{for } x = 1 \text{ billion USD.} \end{cases}$$

Assume that $\pi_X = 0.296$ billion USD and the reinsurance premium is linear $\pi_{\text{re}}(t) = (1 + \theta_{\text{re}})t$, where $\theta_{\text{re}} = 0.6$. Assume that the quota for reinsurance protection $P = 0.05$ billion USD and the initial surplus $V = 0.05$ billion USD. Then the non-randomized policy (22) is of the following form

$$R^*(x) = (x - 0.296)^+ \mathbb{1}_{[0, c^*)}(x).$$

For $c^* = 1$, we have $\pi_{\text{re}}(ER^*) = 0.0008$ billion USD while for any $c^* > 1$, $\pi_{\text{re}}(ER^*) = 0.08528$ billion USD, which is greater than P . Without admitting randomization, the best protection against ruin provides truncated stop loss with $c^* = 1$ billion USD (or with any c^* between 0.3 and 1 billion USD). Then the probability of ruin is 0.075 and the cost of this protection is 0.0008 billion USD. Now take a randomized truncated stop loss such that for the claim $x = 1$ billion USD the reinsurer covers 0.704 billion USD with probability 0.5824. The price of this randomized reinsurance rule is 0.05 billion USD and it decreases the ruin probability from 0.075 to 0.03132.

The same example shows that for discrete claim distributions there might exist a nonrandomized rule, which gives the same ruin probability as truncated stop loss and is less expensive.

Example 5: In Example 4, consider the reinsurance rule with $R \equiv 0$ (i.e., no reinsurance purchase). It gives the probability of ruin 0.075 and costs nothing but is also beaten by the randomized truncated stop loss from Example 4.

Finally, observe that if P is increased from 0.08528 to 0.101882352 billion USD, the insurer could purchase a (nonrandomized) stop loss with the deductible 0.244117648 billion USD, which reduces the probability of ruin to 0.

CONCLUSIONS

To summarize, if the quota P for purchasing reinsurance exceeds $\pi_{\text{re}}(E(X - V - \pi_X + P)^+)$ we have enough money to get a full protection against ruin. The least-expensive full protection reinsurance arrangement then is the stop loss contract with the largest deductible M^* satisfying (6), provided π_{re} is an increasing function of ER . Such deductible M^* exists if π_{re} is positive and continuous on $(0, \infty)$, $\Pr(X > V + \pi_X) > 0$ and, e.g., $\pi_{\text{re}}(EX) < V + \pi_X$. If the first inequality does not hold, the reinsurance is

not necessary. If the second one is reverse, the valuation of the whole portfolio by the reinsurer differs completely from that made by the insurer, which shows that at least one of them uses an improper pricing rule. Thus, typically, an optimal retention M^* should be a unique solution to Equation (12). But even if there are several positive roots of (12), we can take the largest one, which corresponds to the least expensive stop loss reinsurance reducing the probability of ruin to 0.

The insurer's strategy should change radically if P is so small that for all $P' \leq P$ the following inequality holds:

$$\pi_{\text{re}}(E(X - V - \pi_X + P')^+) > P'.$$

Then the least expensive reinsurance policy is the truncated stop loss

$$R^*(t) = \begin{cases} 0 & \text{when } 0 \leq t < V + \pi_X - P, \\ t - V - \pi_X + P & \text{when } V + \pi_X - P \leq t < c^*, \\ 0 & \text{when } t \geq c^*, \end{cases}$$

provided that the total claim distribution is absolutely continuous and

$$\pi_{\text{re}}(t + (1 - F(V + \pi_X - P))(P - \pi_{\text{re}}(t))) \leq P$$

for all $t \in [0, \pi_{\text{re}}^{-1}(P))$ (see Corollary 1). The above condition is related to bounding the marginal safety loading coefficients of the reinsurer by the global safety loading of the insurer plus the insurer's surplus after purchasing the policy expressed in EX (see (20)).

On the other hand, if the above inequality does not hold, truncation stop loss need not be an optimal contract. In "Numerical Examples," an example is given where a more expensive truncated stop loss provides a higher ruin probability than a less expensive one in such a situation. This shows that Corollary 1 is no longer true when this assumption is dropped. It is an intriguing question, whether a similar paradox may also happen for other optimality criteria.

The truncation point c^* is a solution to the equation $ER^* = \pi_{\text{re}}^{-1}(P)$. If the claim distribution is discrete, it may happen that for some $c^* > V + \pi_X - P$

$$ER^* < \pi_{\text{re}}^{-1}(P)$$

and simultaneously $ER' > \pi_{\text{re}}^{-1}(P)$, where

$$R'(t) = \begin{cases} R^*(t), & t \neq c^*, \\ c^* - V - \pi_X + P, & t = c^*. \end{cases}$$

This means that we have more money than necessary to purchase R^* but not enough to purchase R' . Then the randomized rule R_γ , which is equal to R^* for claims different

from c^* and otherwise to $c^* - V - \pi_X + P$, with probability γ , and to 0, with probability $1 - \gamma$, produces a lower ruin probability than R^* and

$$ER_\gamma = \pi_{\text{re}}^{-1}(P)$$

for properly chosen $\gamma \in (0, 1)$. The improvement via randomization never happens if X has absolutely continuous distribution and it vanishes with decreasing the probabilities of atoms. But even if the distribution of X cannot be accurately approximated by an absolutely continuous distribution, the truncated stop loss reinsurance cannot be improved in the probability of survivor by any other nonrandomized policy, if the assumptions of Proposition 2 are satisfied.

REFERENCES

- Amsler, M. H., 1993, Reassurance du risque de ruine, *Bulletin of the Swiss Association of Actuaries* 91(1): 33-39 (see also *Insurance: Mathematics and Economics*, 12(1): 73, 1993).
- Arrow, K., 1951, Alternative Approaches to the Theory of Choice in Risk-Taking Situations, *Econometrica*, 19(4): 404-433.
- Arrow, K., 1963, Uncertainty and the Welfare Economics of Medical Care, *American Economic Review*, 53: 941-973.
- Bühlmann, H., 1970, *Mathematical Methods in Risk Theory* (New York: Springer).
- Daykin, C. D., T. Pentikäinen, and M. Pesonen, 1993, *Practical Risk Theory for Actuaries* (London: Chapman & Hall).
- Gajek, L., and D. Zagrodny, 2000, Insurer's Optimal Reinsurance Strategies, *Insurance: Mathematics and Economics*, 27: 105-112.
- Gajek, L., and D. Zagrodny, 2004, Optimal Reinsurance Under General Risk Measures, *Insurance: Mathematics and Economics*, 34: 227-240.
- Kaluszka M., 2001, Optimal Reinsurance Under Mean-Variance Premium Principles, *Insurance: Mathematics and Economics*, 28: 61-67.
- Kass, R., M. Goovaerts, J. Dhaene, and M. Denuit, 2001, *Modern Actuarial Risk Theory* (Boston: Kluwer Academic Publishers).
- Klugman, S. A., H. H. Panjer, and G. E. Willmot, 1998, *Loss Models: From Data to Decisions* (New York: Wiley).
- Lehmann, E., 1985, *Testing Statistical Hypothesis*, 2nd edition (New York: Wiley).
- Mazur, T., 2000, *Risk Management via Reinsurance*, Diploma thesis, Technical University of Łódź, Łódź (in Polish).
- Pesonen, M. I., 1984, Optimal Reinsurances, *Scandinavian Actuarial Journal*, 84: 65-90.
- Verlaak, R., and J. Beirlant, 2003, Optimal Reinsurance Programs. An Optimal Combination of Several Reinsurance Protections on a Heterogeneous Insurance Portfolio, *Insurance: Mathematics and Economics*, 33: 381-403 (presented on the occasion of Insurance: Mathematics and Economics Congress, Lisbon, 2002).

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.