

The Impact of Regret on the Demand for Insurance

Michael Braun and Alexander Muermann¹

The Wharton School

University of Pennsylvania

September 2003

¹Braun: Department of Operations and Information Management, The Wharton School, University of Pennsylvania, 500 Jon M. Huntsman Hall, Philadelphia, PA 19104-6340, USA (email: mbraun2@wharton.upenn.edu); Muermann: Department of Insurance and Risk Management, The Wharton School, University of Pennsylvania, 3641 Locust Walk, Philadelphia, PA 19104-6218, USA (email: muermann@wharton.upenn.edu). We wish to thank Terry Connolly, David Croson, Rachel Croson, Neil Doherty, Christian Gollier, Paul Kleindorfer, Howard Kunreuther, Olivia Mitchell, Maurice Schweitzer, Kent Smetters, and participants in the Wharton School Applied Economics workshop for their valuable comments.

Abstract

We examine optimal insurance purchase decisions of individuals that exhibit behavior consistent with Regret Theory. Our model incorporates a utility function that assigns a disutility to outcomes that are ex-post suboptimal, and predicts that individuals with regret-theoretical preferences adjust away from the extremes of full insurance and no insurance coverage. This prediction holds for both coinsurance and deductible contracts, and can explain the frequently observed preferences for low deductibles in markets for personal insurance.

1 Introduction

The canonical models of insurance demand treat the insurance purchase decision as an expected utility maximization problem. It is well-established from these models that a risk-averse individual will purchase full insurance when the insurance contract is fairly priced and less than full insurance with a positive proportional loading factor (Mossin 1968; Schlesinger 1981, 2000). This result applies to insurance policies of both a coinsurance structure (where indemnity is a proportion of the loss) and a deductible structure (where indemnity is a payoff for losses above a certain threshold). The supply of insurance in these models is assumed to be exogenous and individuals' preferences are assumed to be consistent with Expected Utility Theory (EUT).¹

Given the plethora of observed deviations from EUT (e.g., Allais Paradox, preference reversals), it is reasonable to expect that observed insurance decisions would not strictly adhere to the predictions of EUT-based models. The robustness of these predictions when the assumption of an EU-maximizing insurance customer is relaxed are mixed. For example, Machina (2000) shows that some predictions of EUT-based insurance demand models remain valid as long as the individual is risk averse; conditions of outcome convexity and linearity in probabilities are unnecessary for both coinsurance and deductible contracts. Doherty and Eeckhoudt (1995) found when applying the "Dual Theory" of Yaari (1987) (where expected utility is not linear in probabilities), that only corner solutions (full insurance or no insurance) can be optimal for coinsurance contracts and that optimal solutions differ from the EUT model (and may be interior) for deductible contracts. Approaches based on the Savage (1954) "minimax regret" rule have also attempted to explain insurance demand.² Razin (1976) found that minimax regret predicted preferences for positive deductibles even when insurance is fairly priced. Briys and Louberge (1985) showed that decision-makers using the Hurwicz criterion (where decision-makers are assumed to assign subjective weights to the best and worst possible outcomes) may prefer either zero or positive deductibles at any loading factor. Clearly, the

¹Practitioners in the insurance industry often define coinsurance as a minimum amount of the total value of a property that must be insured. Our use of the term "coinsurance" is, however, consistent with its usage in the academic literature.

²One of the earliest models of decision making is the "minimax regret" decision rule of Savage (1954), which incorporates preferences to minimize exposure to ex-post suboptimal outcomes. This model defines regret as the difference between the outcome of each decision alternative and the best possible outcome among all alternatives for each ultimate state of the world. A decision-maker applying minimax regret would determine the highest level of regret that could possibly occur for each decision alternative (among all possible outcomes), and then choose the decision alternative with the lowest of these maximum regret levels (Zeelenberg 1999). The minimax regret rule is useful because it places an upper bound on the regret that one would experience. However, it does not take into account the probabilities of any of the states of the world occurring. A decision-maker could potentially pass up an alternative that would offer a near-certain chance of low regret because of a possibility of an occurrence a low probability event that would induce a large amount of regret.

assumptions on the structure of the utility function and decision rules matter when determining the optimal insurance level.

In this paper we propose an alternative, axiomatically supported utility function that explains how a “regret-averse” individual would hedge his own risk differently from an EU-maximizing individual. This function, derived from Regret Theory as proposed by Loomes and Sugden (1982) and Bell (1982) and axiomatized by Sugden (1993) and Quiggin (1994), defines regret as the disutility of not having chosen the ex-post optimal alternative. Behavior of this sort has been observed at various levels of significance in both the laboratory and the field, and we believe that the presence of this behavior may help explain anomalies in the demand for insurance. The major difference between our model and the decision rules discussed in the previous paragraph is that Regret Theory implies that a second attribute—regret—is considered along with wealth as part of the objective function. The previous attempts at this subject assume only one attribute: wealth in the case of EUT, non-EUT and Dual Theory, and gaps between actual and ex-post optimal outcomes in the case of minimax regret. The present approach assumes that both wealth and regret play a part in the decision-making process, as does the probability distribution on possible states of the world.

Our results predict that preferences that are consistent with the Regret Theory axioms of Sugden (1993) and Quiggin (1994) (i.e., preferences that are “regret-theoretical”) should have a “mitigating” effect on insurance decisions when compared to preferences that are consistent with EUT axioms of Savage (1954). We find that when large amounts of insurance are predicted by EUT, a regret-averse individual (i.e., an individual with regret-theoretical preferences) should choose less insurance, and when small amounts of insurance are predicted to be optimal by EUT, the regret-averse individual prefers to purchase more insurance. These results hold for insurance contracts of both the coinsurance form and deductible form, and the direction of the impact of regret depends on the proportional loading factor on the premium of these contracts. We also find that the “critical loading factor” at which the impact of regret goes from positive to negative is independent of the level of regret incorporated in the model. In essence, we show that individuals with regret-theoretical preferences would tend to “hedge their bets,” taking into account the possibility that their decisions may turn out to be ex-post suboptimal. For loading factors that are sufficiently high, these predictions can help explain the preferences for low deductibles that is well-known in the study of insurance.

In the next section, we present a Regret Theoretical Expected Utility (RTEU) function and describe its

properties in light of previous axiomatic and experimental research. In section three, we show how the maximization of RTEU yields optimal insurance choices that are less extreme than those predicted by the conventional expected utility model. In section four, we discuss implications that the overall results have for both the insurance industry and regulatory bodies, including an explanation for observed preferences for low deductibles.

2 Regret-Theoretical Expected Utility

The key assumption of our model is that individuals avoid the unfavorable consequences of experiencing an outcome that is worse than the best that could have been achieved *had the amount of the loss (or lack thereof) been known in advance*. For example, if an individual purchases little insurance, and then incurs a large loss, the individual would experience some additional disutility of not having purchased more insurance ex-ante. One can imagine such an individual “kicking oneself” for not having bought more insurance. Had the individual taken more insurance (in the form of a higher coinsurance rate or lower deductible level), not only would the insurance contract have covered more of the loss, but he would have “felt better” by having made the “right decision.” Among the drivers of this behavior might be the avoidance of emotional regret (“I’ll feel regret if I don’t have enough insurance”) or the need for legitimation (“If I have an accident, I’ll have to explain why I didn’t buy more insurance”). Regardless of whether this disutility is derived from avoiding either a negative emotion or additional negative consequences, we call this disutility “regret” and can reasonably believe a story of “regret avoidance” in the insurance context.

2.1 The RTEU Function

Our model assumes that a representative insurance customer has a two-attribute utility function of the following form

$$u(w) = v(w) - k \cdot g(v(w^{\max}) - v(w)) \quad (1)$$

where $v : \mathbb{R}_+ \rightarrow \mathbb{R}$ is the *traditional* Bernoulli utility (value) function over monetary positions with $v' > 0$ and $v'' < 0$. $g : \mathbb{R}_+ \rightarrow \mathbb{R}$ is the regret function that depends on the difference between the value of the actual final level of wealth w and the value of the ex-post optimal final level of wealth $w^{\max}$. $w^{\max}$ is the wealth that the individual could have received if he had made the optimal choice with respect to the realized

state of nature (the amount of the loss). We assume that $g' > 0$, $g'' > 0$, and $g(0) = 0$. k is the linear weight placed on the regret component of this utility function.³ If $k = 0$, the individual is a *traditional* risk-averse expected utility maximizer. If $k > 0$, then the utility function of the individual includes some compensation for regret and we call the individual *regret-averse*. Any preferences that can be represented by the RTEU function are said to be *regret-theoretical*. Under these assumptions, all individuals that we consider in this analysis are risk-averse (because of the concavity of v), but only those for whom $k > 0$ are also regret-averse.

This utility function is derived from Regret Theory as formulated by both Bell (1982, 1983) and Loomes and Sugden (1982). Each suggested that decision-makers optimize the expected value of a “modified” utility function of the form

$$u(x, y) = v(x) + g(v(x) - v(y)) \quad (2)$$

where $u(\cdot)$, $v(\cdot)$ and $g(\cdot)$ are defined as above, but x and y are defined as two possible outcomes of a lottery. x is the “chosen” outcome—the outcome on which an individual has bet—and y is the “foregone” outcome. Note that $g(\cdot)$ represents the regret or rejoicing that the decision-maker experiences as a result of receiving x versus not receiving y . If it turns out that $x > y$, then the decision-maker made the “correct” choice and gains some additional utility by having passed up the foregone alternative. If $x < y$, then the decision-maker experiences disutility from having forgone the possibility of doing better had he chosen the foregone alternative. Thus, Regret Theory assumes not only that decision-makers experience regret, but also that the anticipation of experiencing regret is factored into the decision-making processes (Larrick and Boles 1995).

Unlike the formulation of Regret Theory of equation (2), which considers two outcomes of a lottery, we are concerned with preferences among actions. It is straightforward to convert the modified utility function from equation (2) into an expected utility function of the form

$$EU(x, y) = \int (v(x_\theta) + k \cdot g(v(x_\theta) - v(y_\theta))) dF(\theta). \quad (3)$$

where $F(\theta)$ is a cumulative distribution function that reflects the subjective beliefs about realizations of

³This utility function is essentially a linear combination of a value for wealth and a value for regret. One could rescale the coefficients on $v(w)$ from 1 to β and on $g(\cdot)$ from k to $(1 - \beta)$, where $\beta = \frac{1}{k+1}$. This transformation places linear weights on two attributes—monetary utility and regret—and allows the decision-maker to maximize the resulting function. In this paper we have chosen to use k instead of β purely for mathematical convenience.

states of the world θ . x_θ is the outcome in state of the world θ that accrues when the chosen action x is taken, and y_θ is the outcome that would have been accrued in state θ had action y been chosen. This notion of expected utility in Regret Theory is more akin to that of Savage (1954) than of von Neumann-Morgenstern, in that we are comparing preferences for actions, assuming subjective probabilities, rather than bets on lotteries given known probabilities. It is clear, however, that equation (3) violates the Axiom of Independence—the payoff of the foregone action affects the value of the chosen one. Thus, the Savage axioms for subjective expected utility cannot be represented by an expected utility function of this form.

We resolve this issue by proposing a Regret-Theoretical Expected Utility (RTEU) function of the form

$$RTEU = \int (v(w_\theta) - k \cdot g(v(w_\theta^{\max}) - v(w_\theta))) dF(\theta) \quad (4)$$

that is consistent with the axioms of Regret Theory of Sugden (1993) and the Axiom of Irrelevance of Statewise Dominated Alternatives (ISDA) proposed by Quiggin (1994). ISDA requires the decision-maker to ignore any actions in the feasible set that are statewise dominated by other actions in the set. The critical consequence of Quiggin’s ISDA is that if a decision-maker’s preferences are consistent with ISDA and the Sugden axioms, then the regret associated with a given action in a particular state of nature depends only on the actual outcome and the best possible outcome that the individual could have attained in that same state of nature. Hence, RTEU as expressed in equation (4) represents preferences that are consistent with these axioms. And since the Sugden axioms are essentially a reformulation of those of Savage (1954), we have a normative basis for RTEU that allows us to use the model to analyze insurance choices. This result also allows us to focus solely on “regret” and its associated disutility, as opposed to earlier formulations of Regret Theory that also allow for “rejoicing” when the “better” outcome is chosen for the eventual state of the world. In fact, because $g(0) = 0$ and because one can never do better than the best possible outcome, we have eliminated “rejoicing” from the regret/rejoice model altogether. We can then restrict $k \geq 0$ as measure of the influence of regret on the decision.

2.2 Behavioral Characteristics

We have already defined regret as the disutility an individual experiences from the value gap between an actual outcome and the best possible outcome that one could have attained in a particular state of nature.

Naturally, the selection of an RTEU-maximizing decision involves some trade-off between maximizing the value from the actual outcome and minimizing regret. Any evidence that individuals behave in a regret-avoiding way is consistent with the RTEU function, and for our purposes, regret-theoretic behavior is any behavior that can be modelled by an expected utility function of the RTEU form. In the psychology literature, however, regret often has specific meaning as an emotion or driver of behavior that we do not consider explicitly in this paper.

Nevertheless, there are some behavioral constructs that are often considered as part of Regret Theory, but are *not*, in our context, consistent with RTEU. First, we emphasize that “regret” is not the same as “disappointment.” Zeelenberg et al. (2000) make the distinction clear. “Regret is assumed to originate from comparisons between the factual decision outcome and a counterfactual outcome that might have been *had one chosen differently*; disappointment is assumed to originate from a comparison between the factual decision outcome and a counterfactual outcome that might have been *had another state of the world occurred*” (p. 529). The distinction between disappointment and regret can be examined in terms of the reference point against which the decisions are made (Loomes 1988). When a decision-maker minimizes regret, the reference point for each action is the outcome that is accrued in the same state of the world from the *other* action. If he minimizes disappointment, each action is measured against the ex-ante expectation of what the payoff could have been for the same action. In the insurance context, we are assuming that the disutility comes from an individual predicting how he would feel after making the wrong decision. The ex-post assessment we are considering is “I should have bought more (or less) insurance” and not “I wish I didn’t incur that loss.” Thus, our interest is clearly in regret and not disappointment.

Second, we note that there is a difference between avoiding regret by making regret-avoiding decisions, and avoiding regret by suppressing regret-inducing information about the outcome of the foregone alternative. Regret-avoiding decisions are of the type we consider in the present work—a subject maximizes his RTEU function. Suppressing regret-inducing information implies that if all signals regarding the eventual state of the world were withheld—and the outcome that would have been experienced had the foregone action been chosen—then the individual could not possibly experience any disutility from that action. Bell (1983), in fact, defines a regret premium as the amount of money a decision-maker is willing to pay to cancel, or suppress the knowledge of the outcome of, a foregone lottery. In our context, we are concerned only with the first approach to avoiding regret: making appropriate ex-ante decisions. It would not be realistic to model a

world in which one does not know the amount of the loss that one incurs.

Finally, we emphasize that regret-avoidance does not necessarily imply either risk-seeking or risk-avoiding behavior. One could be regret-averse and either risk-averse or risk-seeking. The separation of these two concepts is explained by Zeelenberg (1999), who summarizes an experiment (and two subsequent variations) that offer confirmatory evidence that individuals attempt to avoid regret. In the first, subjects are asked to choose between two gambles, one being “risky” and the other being “safe.” In the first condition of the experiment, the outcome of the risky gamble is revealed, and in the second, the outcome of the safe gamble is revealed. In addition, the subject sees that outcome of his chosen gamble. Thus, a subject who chooses the risky gamble will know the outcome only of the risky gamble in the first condition and both gambles in the second, and the subject who chooses the safe gamble will know the outcome of both gambles in the first condition and only the safe gamble in the second. Subjects most frequently chose the gamble that was to be revealed in that condition, demonstrating a preference to not know the outcome of the foregone alternative. Hence, the subjects chose actions that minimize regret, independent of their decisions to play “risky” or “safe.” But in this case, regret is avoided by *suppressing information about the foregone alternative*, and not by choosing regret-minimizing actions. Nevertheless, the implications to the insurance story are clear. One could experience regret by either not purchasing enough insurance that one ultimately needs (risk-averse behavior avoids this condition), or by purchasing more insurance than is ultimately needed (a compensatory reduction in insurance coverage is risk-seeking).

2.3 Empirical Justification of Regret Theory and RTEU

There exists a significant body of evidence to suggest that Regret Theory can explain deviations from EUT that have been observed in both the laboratory and in the field. Loomes and Sugden (1982) show how Kahneman and Tversky’s examples of preference inconsistencies and intransitivities that were explained by Prospect Theory (Kahneman and Tversky 1979) can also be explained by Regret Theory. Bell (1982) illustrates how Regret Theory can help explain simultaneous preferences for insurance and gambling. Experiments to test the specific impact of regret (as opposed to other drivers such as disappointment) on preferences were conducted by Loomes and Sugden (1987) with weak results, but later replication and refinement by Loomes (1988) demonstrated that the influence of regret is, in fact, quite strong. Loomes et al. (1992) subsequently confirmed experimentally that regret is a significant factor when preferences violate stochastic

dominance. And while Starmer and Sugden (1993) showed that effects related to the presentation of the lotteries in the earlier experiments may have strengthened the earlier results, the impact of regret is only marginally less significant than previously thought. Furthermore, these qualifications still do not apply to our particular model where optimal decisions are continuous within contiguous states of the world.

Attempts outside of the laboratory to validate the importance of regret in decision-making are more encouraging. For example, Connolly and Reb (2003), through work on omission bias, found that decisions about whether to get a vaccination tended to be driven by the avoidance of regret. These results are interesting in our context because the anticipated regret works in two directions: some subjects did not get shots because they wanted to avoid risks from the vaccine, while others actively sought out vaccination so they could avoid the disease itself. This decision is similar to the insurance case, in which an individual must balance regret from buying too little insurance with regret from buying too much. Baron and Hershey (1988) did find consistent evidence that the ex-post evaluation of the quality of a decision is influenced by the favorability of the final outcomes. In other words, a decision after the fact is more likely to be considered “good” if the final outcome was good. The authors report a weakly positive change in the ex-post quality of decisions when the attention of subjects was focused on these relative outcomes of the alternatives. However, in one experiment, the authors did find significant evidence that a decision is more likely to be considered good, if the difference between the outcomes of the chosen and foregone alternatives is high. Although this result could constitute “rejoicing” as well as regret, it is still consistent with the idea of anticipating ex-post suboptimal outcomes.

Of course, even if the experimental evidence on regret is inconclusive, it does not mean that regret is not a factor in decision-making (Baron 2000, p. 265). As we develop our model of regret-induced decision making, we will show below that regret can perturb optimal decisions in one direction in some situations, but in other directions for others.

3 The Insurance Model

We now turn our attention to the application of the RTEU function to optimal insurance purchase decisions. Suppose an individual is endowed with an initial level of wealth $w_0 \geq 0$. With probability $1 - q$, this individual faces a monetary, random loss X which is characterized by a cumulative distribution function

$F : [0, w_0] \rightarrow \mathbb{R}$ with $F(0) = 0$ and $F(w_0) = 1$. This probability distribution is conditional on there being a loss and all subsequent expected value operators $E[\cdot]$ are taken with respect to this conditional distribution. An insurance company offers a menu of indemnity insurance contracts to the individual for premiums that are set equal to the expected indemnity plus proportional loading factor $\lambda \geq 0$. This pricing behavior can be justified economically by assuming that the insurer is risk-neutral in a perfectly competitive insurance market with proportional transaction costs but no entry costs. We exclude any informational asymmetries that would give rise to moral hazard or adverse selection problems. The individual's decision making is assumed to be representable by RTEU maximization subject to a Bernoulli utility function $u : \mathbb{R}_+ \rightarrow \mathbb{R}$, which, due to regret-aversion, exhibits the following structure:

$$u(w) = v(w) - k \cdot g(v(w^{\max}) - v(w))$$

for all levels of final wealth $w \geq 0$. (The definitions and conditions for these functions were presented in section 2.1).

In the next two subsections, we examine how regret impacts the individual's insurance purchasing decision if the insurer offers a coinsurance and a deductible policy. In all cases, when we refer to the amount of insurance, we refer to the coinsurance rate or deductible level and not an upper limit on coverage.

3.1 Coinsurance Policy

Suppose an insurer offers a coinsurance contract with coinsurance rates $\alpha \in [0, 1]$. The indemnity schedule $I : [0, w_0] \rightarrow \mathbb{R}$ is given by

$$I(x) = \alpha x$$

for all realized losses $x \in \mathbb{R}_+$ of the random variable X and premiums are equal to

$$\begin{aligned} P(\alpha) &= (1 + \lambda)(1 - q)E[I(X)] \\ &= (1 + \lambda)(1 - q)\alpha E[X], \end{aligned} \tag{5}$$

where $\lambda \geq 0$ is the proportional loading factor.

To determine the optimal insurance purchasing decision we first deduce the ex-post optimal level of final

wealth $w^{\max}$.

Lemma 1 *The ex-post optimal insurance purchasing decision is full insurance ($\alpha = 1$) if the realized loss x exceeds $(1 + \lambda)(1 - q)E[X]$ and no insurance ($\alpha = 0$) if there is either no loss or a realized loss x that falls below $(1 + \lambda)(1 - q)E[X]$. The ex-post optimal level of final wealth is $w_L^{\max} = w_0 - \min(x, (1 + \lambda)(1 - q)E[X])$ in case of a loss and $w_{NL}^{\max} = w_0$ in case of no loss.*

Proof. in appendix ■

These ex-post choices are interpreted as the amount of insurance the individual would have purchased had the actual amount of the loss been known in advance. If the loss were less than the premium for full insurance, then the individual would have done better by paying for the loss *instead* of the premium. Similarly, if the amount of the loss were greater than the premium for full insurance, the individual is better off purchasing full coverage for that loss.

Mossin (1968) has shown that a risk-averse individual who is not regret-averse would buy full insurance ($\alpha^* = 1$) if the contract were fairly priced. Partial insurance may be indicated either by a positive loading factor ($\lambda > 0$), or by moral hazard (Holmstrom 1979) or adverse selection (Rothschild and Stiglitz 1976). We now establish another reason why individuals would demand partial insurance in the absence of such information asymmetries.

Proposition 2 *If a coinsurance contract is offered, a regret-averse individual purchases partial insurance ($\alpha^* < 1$) even at a fair price ($\lambda = 0$).*

Proof. in appendix ■

This implies that for a fairly priced contract, there is some regret incurred from purchasing insurance that is not ultimately needed. Consider a regret-averse individual who purchases full insurance under a coinsurance contract that is fairly priced ($\lambda = 0$). If $x > (1 - q)E[X]$, he experiences no regret—he has made the ex-post optimal decision. Otherwise, he experiences a lot of regret from buying any insurance at all. By reducing the coinsurance rate, the individual experiences regret in all states of the world—more regret for $x > (1 - q)E[X]$ and less regret otherwise. Because g is convex, this adjustment reduces the expected disutility of regret. Therefore, full insurance cannot be optimal.

Let $\alpha_k^*(\lambda)$ denote the optimal coinsurance rate for a given loading factor $\lambda \geq 0$ and regret parameter $k \geq 0$. We have already shown that for $k > 0$, $\alpha_k^*(\lambda) < 1$ for all $\lambda \geq 0$. In the following proposition, we

show how the optimal coinsurance level $\alpha_k^*(0)$ responds to changes in regret-aversion, measured by k , if the contract is fairly priced.

Proposition 3 *If a fairly priced coinsurance contract is offered, an individual who is more regret-averse with respect to k purchases less insurance coverage than the individual who is less regret-averse (i.e., $\frac{d\alpha_k^*(0)}{dk} < 0$).*

Proof. in appendix ■

We know from Mossin (1968) that there exists some loading factor $\bar{\lambda} > 0$ above which the individual will purchase no insurance (i.e. $\alpha_0^*(\bar{\lambda}) = 0$). In the next proposition, we show at this $\bar{\lambda}$ that a regret-averse individual purchases some insurance and that more regret-aversion leads to a higher level of insurance.

Proposition 4 *If a coinsurance contract is offered with a loading factor*

$$\bar{\lambda} = \frac{\text{Cov}[v'(w_0 - X), X] + qE[X](E[v'(w_0 - X)] - v'(w_0))}{E[X](qv'(w_0) + (1 - q)E[v'(w_0 - X)])} > 0,$$

then a regret-averse individual ($k > 0$) will purchase some non-zero amount of insurance ($\alpha_k^(\bar{\lambda}) > 0$) while a non-regret-averse individual ($k = 0$) would not buy any insurance ($\alpha_0^*(\bar{\lambda}) = 0$). If the amount of regret goes up (i.e., k increases), then the optimal coinsurance rate increases (i.e., $\frac{d\alpha_k^*(\bar{\lambda})}{dk} > 0$).*

Proof. in appendix ■

The intuition behind this result parallels that of the fairly priced contract. The expected disutility of regret can be reduced by purchasing some amount of insurance.

So far we have derived the following results:

$$\begin{aligned} \alpha_k^*(\lambda) &< 1, \text{ for all } k > 0, \lambda \geq 0 \\ \frac{d\alpha_k^*(0)}{dk} &< 0, \text{ for all } k > 0 \\ \alpha_k^*(\bar{\lambda}) &> 0, \text{ for all } k > 0 \\ \frac{d\alpha_k^*(\bar{\lambda})}{dk} &> 0, \text{ for all } k > 0. \end{aligned}$$

These results mean that for $\lambda = 0$, regret induces an individual to buy partial insurance instead of full insurance, and at $\lambda = \bar{\lambda}$, regret induces an individual to purchase partial insurance instead of no insurance.

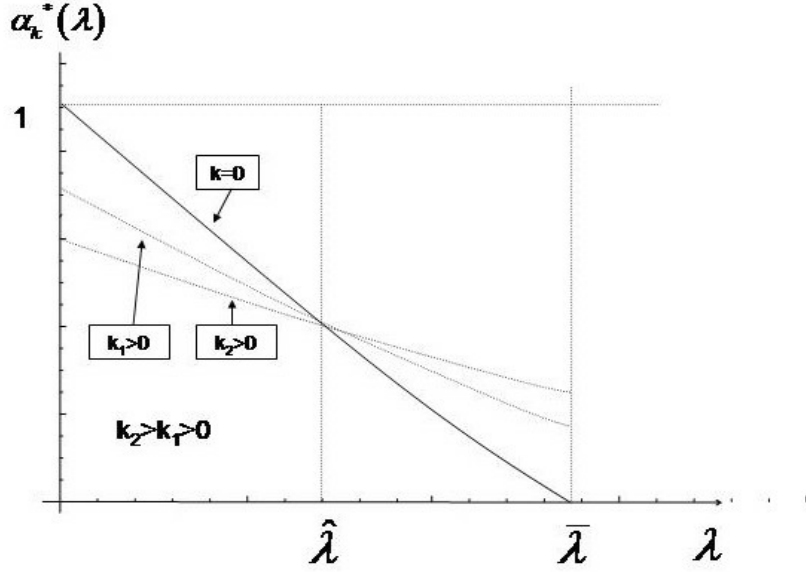


Figure 1: Plots the optimal coinsurance rate against loading factors. For $k = 0$, an individual demands full insurance for $\lambda = 0$ and no insurance for $\lambda = \bar{\lambda}$. For $k > 0$, the individual demands partial insurance at both $\lambda = 0$ and $\lambda = \bar{\lambda}$. The amount of the adjustment increases as k increases.

Finally, there exists a loading factor for which regret has no effect on the insurance choice and this loading factor reflects the point at which the impact of regret switches from less insurance to more insurance.

Proposition 5 *There exists a loading factor $\hat{\lambda} \in (0, \bar{\lambda})$ such that $\alpha_k^*(\hat{\lambda}) = \alpha_0^*(\hat{\lambda})$ for all $k \geq 0$.*

Proof. in appendix ■

We illustrate the effects of regret on the optimal coinsurance rate in Figure 1.

Aside from the endpoints of the $\alpha_0^*(\lambda)$ plot, we do not know the exact *shape* of either curve in the interior. In EUT-based insurance demand models, the slope and concavity of $\alpha_0^*(\lambda)$ are driven by wealth and substitution effects, and this characteristic remains for the RTEU model. However, we see that for low loading factors, adding regret induces the purchase of less insurance, while high loading factors imply that regret leads to purchasing more insurance. These results do *not* imply necessarily that individuals purchase more insurance at higher loading factors. Rather, the perturbations in the optimal insurance demand should be thought of as the amount of insurance demanded by the regret-averse individual *relative* to that of the straight expected utility maximizer. Nor could these results be replicated by altering the risk-aversion of a

single attribute expected utility function, since the EUT predictions at the endpoints (full insurance or no insurance) hold for a risk-averse individual with any concave utility function.

3.2 Deductible Policy

We now turn to insurance policies with deductibles. Suppose an insurer offered a set of deductible contracts with deductible levels $D \in [0, w_0]$. The indemnity schedule is thus

$$\begin{aligned} I(x) &= \max(x - D, 0) \\ &= (x - D)^+ \end{aligned}$$

and the premium is given by

$$P(D) = (1 + \lambda)(1 - q) E[(X - D)^+].$$

Differentiation yields

$$P'(D) = -(1 + \lambda)(1 - q)(1 - F(D)). \quad (6)$$

The individual chooses a deductible level to maximize his expected utility of final wealth. If a loss occurs, this amount is

$$\begin{aligned} w_L(D) &= w_0 - (1 + \lambda)(1 - q) E[(X - D)^+] - X + (X - D)^+ \\ &= w_0 - (1 + \lambda)(1 - q) E[(X - D)^+] - \min(X, D) \end{aligned}$$

and if no loss occurs, this amount is

$$w_{NL}(D) = w_0 - (1 + \lambda)(1 - q) E[(X - D)^+].$$

In the following Lemma, we derive the ex-post optimal final level of wealth, $w^{\max}$.

Lemma 6 *The ex-post optimal insurance purchasing decision if there is no loss is no insurance, and $w_{NL}^{\max} = w_0$. If the loading factor $\lambda \leq \frac{q}{1-q}$, then it is ex-post optimal to buy no insurance if the realized loss $x < (1 + \lambda)(1 - q) E[X]$, and to buy full insurance otherwise. The ex-post optimal level of final wealth $w_L^{\max}$ is*

then given by

$$w_L^{\max} = w_0 - \min(x, (1 + \lambda)(1 - q) E[X])$$

If the loading factor $\lambda > \frac{q}{1-q}$, it is ex-post optimal to buy a deductible level $\bar{D}(\lambda)$ if the realized loss $x \geq (1 + \lambda)(1 - q) E[(X - \bar{D}(\lambda))^+] + \bar{D}(\lambda)$, and to buy no insurance otherwise. $\bar{D}(\lambda)$ is the $\frac{\lambda(1-q)-q}{(1+\lambda)(1-q)}$ th quantile of the loss distribution function F , i.e. $\bar{D}(\lambda) = F^{-1}\left(\frac{\lambda(1-q)-q}{(1+\lambda)(1-q)}\right)$. The ex-post optimal level of final wealth $w_L^{\max}$ is then given by

$$w_L^{\max} = w_0 - \min\left(x, (1 + \lambda)(1 - q) E[(X - \bar{D}(\lambda))^+] + \bar{D}(\lambda)\right)$$

Proof. in appendix ■

The following Proposition shows that an individual who is both risk-averse and regret-averse will demand partial insurance under a deductible policy, even if the contract is fairly priced ($\lambda = 0$).

Proposition 7 *If a deductible insurance contract is offered, a regret-averse individual purchases less than full insurance ($D^* > 0$) even at a fair price ($\lambda = 0$).*

Proof. in appendix ■

Now let $D_k^*(\lambda)$ denote the optimal deductible level for a given loading factor $\lambda \geq 0$ and regret parameter $k \geq 0$. We have already shown that for $k > 0$, $D_k^*(\lambda) > 0$ for all $\lambda \geq 0$. Analogous to the coinsurance case, we will show how the optimal deductible level $D_k^*(0)$ responds to changes in the regret parameter k if the contract is fairly priced.

Proposition 8 *If a fairly priced ($\lambda = 0$) deductible contract is offered, an individual who is more regret-averse with respect to k purchases less insurance coverage than the individual who is less regret-averse (i.e., $\frac{dD_k^*(0)}{dk} > 0$).*

Proof. in appendix ■

Next, we note that there exists some loading factor $\bar{\lambda}$ above which an individual that is risk-averse, but not regret-averse, will purchase no insurance. However, a regret-averse individual will purchase some insurance priced with the loading factor $\bar{\lambda}$ and more regret-aversion leads to a higher level of insurance, i.e. to a lower deductible level.

Proposition 9 *If a deductible contract is offered with a loading factor*

$$\bar{\lambda} = \frac{v'(0)}{qv'(w_0) + (1-q)E[v'(w_0 - X)]} - 1 > 0,$$

then a regret-averse individual ($k > 0$) purchases some non-zero amount of insurance ($D_k^(\bar{\lambda}) < w_0$) while a non-regret-averse individual ($k = 0$) will purchase no insurance ($D_0^*(\bar{\lambda}) = w_0$). If the amount of regret goes up (i.e., k increases), then the optimal deductible level decreases (i.e., $\frac{dD_k^*(\bar{\lambda})}{dk} < 0$).*

We have thus derived results that are analogous to coinsurance, namely

$$\begin{aligned} D_k^*(\lambda) &> 0, \text{ for all } k > 0, \lambda \geq 0 \\ \frac{dD_k^*(0)}{dk} &> 0, \text{ for all } k > 0 \\ D_k^*(\bar{\lambda}) &< w_0, \text{ for all } k > 0 \\ \frac{dD_k^*(\bar{\lambda})}{dk} &< 0, \text{ for all } k > 0. \end{aligned}$$

Finally, we show that there exists a loading factor for which regret has no effect on the insurance choice and that this loading factor is the point at which the impact of regret switches from less insurance to more insurance.

Proposition 10 *There exists a loading factor $\hat{\lambda} \in (0, \bar{\lambda})$ such that $D_k^*(\hat{\lambda}) = D_0^*(\hat{\lambda})$ for all $k \geq 0$.*

Proof. in appendix ■

Again, the loading factor at which the effect of regret switches from “more insurance” to “less insurance” is the same for all levels of regret. Figure 2 illustrates the impact of regret on the optimal deductible level. This impact exactly parallels that of the coinsurance case.

4 Discussion

4.1 Preferences for Low Deductibles

Our results establish conditions under which regret can help explain observed preferences for low deductibles; the regret-averse individual will buy more insurance than the non-regret-averse individual as long as the

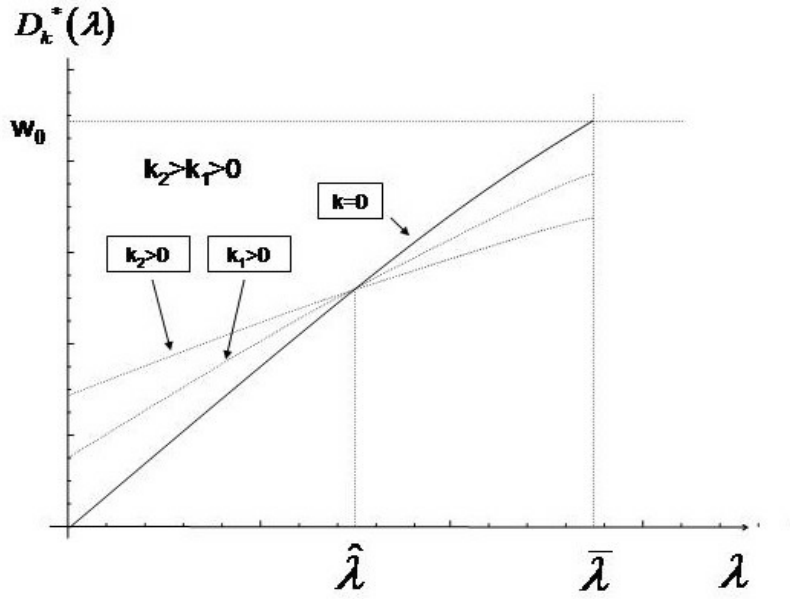


Figure 2: Plots the optimal deductible level against loading factors. For $k = 0$, an individual demands full insurance for $\lambda = 0$ and no insurance for $\lambda = \bar{\lambda}$. For $k > 0$, the individual demands partial insurance at both $\lambda = 0$ and $\lambda = \bar{\lambda}$. The amount of the adjustment increases as k increases.

loading factor is sufficiently high. The existence for the preference for low deductibles is well-known. One of the first empirical studies of the low deductible phenomenon was conducted by Pashigian et al. (1966), who derived ranges into which the insurance premia of an expected utility maximizing individual (assuming a quadratic utility function) must fall. Their results showed that out of a sample of more than 4.8 million insured drivers in 1962, about 53.8 percent chose the lowest deductible and another 45.7 percent chose the next lowest. Yet, none of the premia paid by those drivers fell within the range that would be consistent with maximizing expected utility; only the sliver choosing the highest deductibles fell within that range. In the late 1970's, a proposal in Pennsylvania to place a minimum deductible of \$100 on automobile insurance policies was ultimately rescinded after public outcry, even though such legislation could have saved consumers millions of dollars each year (Cummins, et. al, 1978; Kunreuther 2000) . More recently, Grace et al. (2003) summarized preferences for low deductibles for homeowners' insurance in New York and Florida for catastrophic losses. Preferences for low deductibles have also been indicated experimentally. Johnson et al. (1993) asked subjects to choose between two hypothetical insurance policies. The first policy was structured with a deductible. The second policy was actuarially identical to the first, but had no deductible, a higher premium and a rebate to the customer in the event no claim was filed. More than two-thirds of the subjects preferred the rebate option to the deductible option, suggesting a disproportionate aversion to covering the deductible portion of the total loss. Slovic et al. (1977) conducted experiments of both the "ball in urn" and "simulated business insurance situation" variety that suggest that individuals are more likely to insure against high-probability, low-impact events than against low-probability, high-impact ones of common expected value. These two results demonstrate the preference for individuals to pay in advance in order to avoid future losses, since policies with lower deductibles or higher coinsurance rates offer higher amounts of reimbursement once those losses are accrued.

In considering the impact of our results, we did not consider another possible explanation for preferences for low deductibles. We may observe these preferences because a large number of insurance customers choose the lowest deductible contracts *that are presented to them*. It is possible that we see clustering at the low end of the scale, because the highest deductible contract is just too high.

4.2 Conclusion and Future Research

We have shown that individuals with regret-theoretical preferences will “hedge their bets” when making insurance decisions for coinsurance contracts, with equivalent results for deductible contracts. Regardless of the loading factor, the anticipation of disutility from regret creates a mitigating effect on the insurance purchase decision. When EUT predicts a high optimal level of insurance, the RTEU model compensates for the fact that there will be some states of the world for which no insurance is optimal ex-post. Similarly, if EUT were to predict a low level of insurance, the present model predicts a higher level to take into account those states of the world for which more insurance is optimal ex-post. No matter what insurance decision is made, the individual is adjusting for the possibility of either buying too little insurance if a large loss is incurred, or too much insurance that is never used. Thus, anticipation of regret should prevent an individual from making extreme decisions. These results cannot be explained by risk-aversion alone, since any risk-averse individual will buy full insurance when the insurance is fairly priced (Mossin 1968; Schlesinger 1981, 2000).

These findings have significant implications for sellers of insurance and for regulators. One possible path for future research is to examine the structure of the optimal insurance policy when customers have regret-theoretical preferences. These preferences could induce an insurance company to adjust loading factors or the menu of insurance levels to compensate for the mitigation effect. A regulator, however, may be concerned with incentives for insurance companies to manipulate regret through marketing activities. If marketers could induce an individual's preferences to be more regret-theoretical and, as a result of such a shift the optimal loading factor goes up, then the marketer is essentially raising the price of and the demand for insurance simultaneously. But if the loading factor were too low, inducing regret might lead individuals to purchase too little insurance, placing the customers at a greater risk of experiencing an unmanageable loss. Our present model opens the door for additional experimental and empirical work to better understand how tolerances of risk and regret influence actual insurance decisions. We also see the potential for expanding this analytical approach to the study of regret to areas outside of insurance. Any situation in which individuals “hedge their bets” as a risk-management strategy would be targetable. If we did not continue this research, we would certainly regret it.

A Appendix

A.1 Proof of Lemma 1

If no loss has occurred, then it is ex-post optimal for the individual to have bought no insurance, (i.e. $w_{NL}^{\max} = w_0$). If a loss of severity x occurs, then the individual's final level of wealth is

$$\begin{aligned} w_L(\alpha) &= w_0 - (1 + \lambda)(1 - q)\alpha E[X] - x + \alpha x \\ &= w_0 - x - \alpha[(1 + \lambda)(1 - q)E(X) - x]. \end{aligned}$$

If $x < (1 + \lambda)(1 - q)E[X]$, then $w_L(\alpha)$ is maximized at $\alpha = 0$. Otherwise, $w_L(\alpha)$ is maximized at $\alpha = 1$. Therefore, the optimal ex-post level of final wealth $w_L^{\max}$ for a realized loss of severity x is given by

$$w_L^{\max} = \begin{cases} w_0 - x, & \text{if } x < (1 + \lambda)(1 - q)E[X] \\ w_0 - (1 + \lambda)(1 - q)E[X], & \text{if } x \geq (1 + \lambda)(1 - q)E[X] \end{cases}.$$

A.2 Proof of Proposition 2

The individual's optimization problem is given by

$$\begin{aligned} \max_{\alpha \in [0,1]} E[u(w(\alpha))] &= q(v(w_{NL}(\alpha)) - kg(v(w_0) - v(w_{NL}(\alpha)))) \\ &\quad + (1 - q)(E[v(w_L(\alpha))] - kE[g(v(w_L^{\max}) - v(w_L(\alpha)))) \end{aligned}$$

where

$$\begin{aligned} w_{NL}(\alpha) &= w_0 - (1 + \lambda)(1 - q)\alpha E[X] \\ w_L(\alpha) &= w_0 - (1 + \lambda)(1 - q)\alpha E[X] - X + \alpha X \\ w_L^{\max} &= w_0 - \min(X, (1 + \lambda)(1 - q)E[X]). \end{aligned}$$

The first-order condition (FOC) for this maximization problem is

$$\begin{aligned} &\frac{dE[u(w(\alpha))]}{d\alpha} \\ &= -(1 + \lambda)q(1 - q)v'(w_{NL}(\alpha))E[X](1 + kg'(v(w_0) - v(w_{NL}(\alpha)))) \\ &\quad + (1 - q)E[v'(w_L(\alpha))(X - (1 + \lambda)(1 - q)E[X])(1 + kg'(v(w_L^{\max}) - v(w_L(\alpha))))] \\ &= 0. \end{aligned} \tag{7}$$

The second-order condition (SOC) is

$$\begin{aligned} &\frac{d^2E[u(w(\alpha))]}{d\alpha^2} \\ &= q((1 + \lambda)(1 - q)E[X])^2[v''(w_{NL}(\alpha))(1 + kg'(v(w_0) - v(w_{NL}(\alpha)))) \\ &\quad - kv'^2(w_{NL}(\alpha))g''(v(w_0) - v(w_{NL}(\alpha)))] \\ &\quad + (1 - q)E[(X - (1 + \lambda)(1 - q)E[X])^2(v''(w_L(\alpha))(1 + kg'(v(w_L^{\max}) - v(w_L(\alpha)))) \\ &\quad - kv'^2(w_L(\alpha))g''(v(w_L^{\max}) - v(w_L(\alpha))))]. \end{aligned} \tag{8}$$

Since $g' > 0$ and $g'' > 0$, equation (8) is negative and expected utility a concave function in α . Any solution α^* of the FOC in equation (7) is thus a global maximum. Evaluating the FOC at the full insurance point $\alpha = 1$ yields

$$\begin{aligned} & \frac{dE[u(w(\alpha))]}{d\alpha} \Big|_{\alpha=1} \\ = & -(1+\lambda)q(1-q)v'(w_{NL}(1))E[X](1+kg'(v(w_0)-v(w_{NL}(1)))) \\ & + (1-q)v'(w_L(1))E[(X-(1+\lambda)(1-q)E[X])(1+kg'(v(w_L^{\max})-v(w_L(1))))]. \end{aligned} \quad (9)$$

For no loss, we have $g'(v(w_0)-v(w_{NL}(1))) > g'(0)$ as g is increasing and convex. For loss realizations $x \geq (1+\lambda)(1-q)E[X]$, $w_L^{\max} = w_L(1)$, so

$$g'(v(w_L^{\max})-v(w_L(1)))(x-(1+\lambda)(1-q)E[X]) = g'(0)(x-(1+\lambda)(1-q)E[X]). \quad (10)$$

For realizations $x < (1+\lambda)(1-q)E[X]$, $w_L^{\max} = w_0 - x > w_L(1)$. Since g and v are both increasing and g is convex, we get

$$g'(v(w_L^{\max})-v(w_L(1)))(x-(1+\lambda)(1-q)E[X]) < g'(0)(x-(1+\lambda)(1-q)E[X]). \quad (11)$$

If we match the terms in (9) pairwise with the terms in (11) and (10), then for $k > 0$, we see that

$$\begin{aligned} & \frac{dE[u(w(\alpha))]}{d\alpha} \Big|_{\alpha=1} \\ < & -(1+\lambda)q(1-q)v'(w_{NL}(1))E[X](1+kg'(0)) - (1-q)v'(w_L(1))E[X](1+kg'(0))(\lambda(1-q)-q) \\ = & -\lambda(1-q)v'(w_0-(1+\lambda)(1-q)E[X])E[X](1+kg'(0)) \\ \leq & 0 \end{aligned}$$

for all $\lambda \geq 0$. This implies that at the full insurance level, the individual can increase expected utility by reducing the coinsurance rate. Therefore, the individual chooses a coinsurance rate $\alpha^* < 1$ even if the contract is actuarially fairly price ($\lambda = 0$).

A.3 Proof of Proposition 3

We need to show that $\frac{d\alpha_k^*(0)}{dk} < 0$. Applying the total differential to the FOC (equation (7)) at the optimal $\alpha_k^*(\lambda)$ leads to

$$\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha^2} \Big|_{\alpha=\alpha_k^*(\lambda)} \cdot d\alpha_k^*(\lambda) + \frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(\lambda)} \cdot dk = 0.$$

Hence

$$\frac{d\alpha_k^*(\lambda)}{dk} = - \frac{\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(\lambda)}}{\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha^2} \Big|_{\alpha=\alpha_k^*(\lambda)}}. \quad (12)$$

From the SOC (equation 8) we know that $\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha^2} \Big|_{\alpha=\alpha_k^*(\lambda)} < 0$. Therefore,

$$\text{sign} \left(\frac{d\alpha_k^*(\lambda)}{dk} \right) = \text{sign} \left(\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(\lambda)} \right).$$

The cross-partial derivative at $\lambda = 0$ is given by

$$\begin{aligned} & \frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(0)} \\ &= -q(1-q)v'(w_{NL}(\alpha_k^*(0)))E[X]g'(v(w_0) - v(w_{NL}(\alpha_k^*(0)))) \\ & \quad + (1-q)E[v'(w_L(\alpha_k^*(0)))(X - (1-q)E[X])g'(v(w_L^{\max}) - v(w_L(\alpha_k^*(0))))]. \end{aligned}$$

From both the FOC (equation (7)) and the identity $Cov(Y, Z) = E(YZ) - E(Y)E(Z)$, we deduce that for $k > 0$

$$\begin{aligned} & \frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(0)} \tag{13} \\ &= \frac{1}{k} [q(1-q)v'(w_{NL}(\alpha_k^*(0)))E[X] - (1-q)E[v'(w_L(\alpha_k^*(0)))(X - (1-q)E[X])] \\ &= \frac{1}{k} (1-q) [qv'(w_{NL}(\alpha_k^*(0)))E[X] - Cov(v'(w_L(\alpha_k^*(0))), X) - qE[v'(w_L(\alpha_k^*(0)))]E[X]] \\ &= -\frac{1}{k} (1-q) [qE[X](E[v'(w_L(\alpha_k^*(0))]) - v'(w_{NL}(\alpha_k^*(0)))) + Cov(v'(w_L(\alpha_k^*(0))), X)]. \end{aligned}$$

The concavity of v implies that $Cov(v'(w_L(\alpha_k^*(0))), X) > 0$ and $E[v'(w_L(\alpha_k^*(0)))] > v'(w_{NL}(\alpha_k^*(0)))$. Therefore $\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(0)} < 0$ and we conclude that

$$\text{sign} \left(\frac{d\alpha_k^*(0)}{dk} \right) = \text{sign} \left(\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(0)} \right) < 0.$$

So for fairly priced insurance, more regret (with respect to k) induces an individual to buy less insurance.

A.4 Proof of Proposition 4

For $k = 0$,

$$\begin{aligned} & \frac{dE[u(w(\alpha))]}{d\alpha} \Big|_{\alpha=0} \\ &= -(1+\lambda)q(1-q)v'(w_0)E[X] + (1-q)E[v'(w_0 - X)(X - (1+\lambda)(1-q)E[X])]. \end{aligned}$$

We find the loading factor $\bar{\lambda}$, at which a risk-averse individual who does not consider regret buys no insurance coverage, by solving

$$-(1+\bar{\lambda})q(1-q)v'(w_0)E[X] + (1-q)E[v'(w_0 - X)(X - (1+\bar{\lambda})(1-q)E[X])] = 0, \tag{14}$$

which yields

$$\bar{\lambda} = \frac{Cov[v'(w_0 - X), X] + qE[X](E[v'(w_0 - X)] - v'(w_0))}{E[X](qv'(w_0) + (1-q)E[v'(w_0 - X)])}.$$

For $k > 0$ and $\lambda = \bar{\lambda}$, we have

$$\begin{aligned} & \frac{dE[u(w(\alpha))]}{d\alpha} \Big|_{\alpha=0} \\ &= -(1+\bar{\lambda})q(1-q)v'(w_0)E[X](1+kg'(0)) \\ & \quad + (1-q)E[v'(w_0 - X)(X - (1+\bar{\lambda})(1-q)E[X])(1+kg'(v(w_L^{\max}) - v(w_0 - X)))]. \end{aligned}$$

Using equation (14) we find that

$$\begin{aligned} & \frac{dE[u(w(\alpha))]}{d\alpha} \Big|_{\alpha=0} \\ &= - (1 + \bar{\lambda}) q (1 - q) v'(w_0) E[X] k g'(0) \\ & \quad + (1 - q) k E[v'(w_0 - X) (X - (1 + \bar{\lambda}) (1 - q) E[X]) g'(v(w_L^{\max}) - v(w_0 - X))] . \end{aligned} \quad (15)$$

For realizations $x > (1 + \bar{\lambda}) (1 - q) E[X]$, $w_L^{\max} > w_0 - x$. Therefore, since g and v are both increasing and g is a convex function,

$$g'(v(w_L^{\max}) - v(w_0 - x)) (x - (1 + \bar{\lambda}) (1 - q) E[X]) > g'(0) (x - (1 + \bar{\lambda}) (1 - q) E[X]) . \quad (16)$$

For realizations $x \leq (1 + \bar{\lambda}) (1 - q) E[X]$, $w_L^{\max} = w_0 - x$ so

$$g'(v(w_L^{\max}) - v(w_0 - x)) (x - (1 + \bar{\lambda}) (1 - q) E[X]) = g'(0) (x - (1 + \bar{\lambda}) (1 - q) E[X]) . \quad (17)$$

By matching equations (16) and (17) to equation (15) and by applying equation (14) we get

$$\begin{aligned} & \frac{dE[u(W(\alpha))]}{d\alpha} \Big|_{\alpha=0} \\ &> - (1 + \bar{\lambda}) q (1 - q) v'(w_0) E[X] k g'(0) + (1 - q) k g'(0) E[v'(w_0 - X) (X - (1 + \bar{\lambda}) (1 - q) E[X])] \\ &= k g'(0) [- (1 + \bar{\lambda}) q (1 - q) v'(w_0) E[X] + (1 - q) E[v'(w_0 - X) (X - (1 + \bar{\lambda}) (1 - q) E[X])]] \\ &= 0 . \end{aligned}$$

Since $E[u(w(\alpha))]$ is a concave function in α , we can conclude that $\alpha_k^*(\bar{\lambda}) > 0$.

According to (13), for $k > 0$ and $\lambda = \bar{\lambda}$

$$\begin{aligned} & \frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(\bar{\lambda})} \\ &= \frac{1}{k} [(1 + \bar{\lambda}) q (1 - q) v'(w_{NL}(\alpha_k^*(\bar{\lambda}))) E[X] - (1 - q) E[v'(w_L(\alpha_k^*(\bar{\lambda}))) (X - (1 + \bar{\lambda}) (1 - q) E[X])]] . \end{aligned}$$

For realizations $x > (1 + \bar{\lambda}) (1 - q) E[X]$, $w_L(\alpha_k^*(\bar{\lambda})) > w_0 - x$ and thus

$$v'(w_L(\alpha_k^*(\bar{\lambda}))) (x - (1 + \bar{\lambda}) (1 - q) E[X]) < v'(w_0 - x) (x - (1 + \bar{\lambda}) (1 - q) E[X])$$

as v is an increasing, concave function. For $x \leq (1 + \bar{\lambda}) (1 - q) E[X]$, $w_L(\alpha_k^*(\bar{\lambda})) < w_0 - x$ and thus

$$v'(w_L(\alpha_k^*(\bar{\lambda}))) (x - (1 + \bar{\lambda}) (1 - q) E[X]) \leq v'(w_0 - x) (x - (1 + \bar{\lambda}) (1 - q) E[X]) .$$

This leads to

$$E[v'(w_L(\alpha_k^*(\bar{\lambda}))) (X - (1 + \bar{\lambda}) (1 - q) E[X])] < E[v'(w_0 - X) (X - (1 + \bar{\lambda}) (1 - q) E[X])]$$

and thus

$$\begin{aligned}
& \frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(\bar{\lambda})} \\
& > \frac{1}{k} [(1+\bar{\lambda})q(1-q)v'(w_{NL}(\alpha_k^*(\bar{\lambda})))E[X] - (1-q)E[v'(w_0-X)(X-(1+\bar{\lambda})(1-q)E[X]))] \\
& = \frac{1}{k} [(1+\bar{\lambda})q(1-q)v'(w_{NL}(\alpha_k^*(\bar{\lambda})))E[X] + (1+\bar{\lambda})q(1-q)v'(w_0)E[X]]
\end{aligned}$$

because of (14). We then conclude that

$$\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(\bar{\lambda})} > \frac{1}{k} (1+\bar{\lambda})q(1-q)E[X](v'(w_{NL}(\alpha_k^*(\bar{\lambda}))) + v'(w_0)) > 0$$

as v is an increasing function. As a result,

$$\text{sign}\left(\frac{d\alpha_k^*(\bar{\lambda})}{dk}\right) = \text{sign}\left(\frac{\partial^2 E[u(w(\alpha))]}{\partial \alpha \partial k} \Big|_{\alpha=\alpha_k^*(\bar{\lambda})}\right) > 0.$$

A.5 Proof of Proposition 5

For an arbitrarily fixed $k > 0$, $\alpha_0^*(0) = 1 > \alpha_k^*(0)$ and $\alpha_0^*(\bar{\lambda}) = 0 < \alpha_k^*(\bar{\lambda})$. Since all functions are sufficiently smooth and $\alpha_k^*(\lambda)$ is a continuous function in λ , there must thus exist a $\hat{\lambda}_k$ with $0 < \hat{\lambda}_k < \bar{\lambda}$ such that

$$\alpha_k^*(\hat{\lambda}_k) = \alpha_0^*(\hat{\lambda}_k).$$

As $\alpha_k^*(\hat{\lambda}_k)$ is an inner solution, $\hat{\lambda}_k$ is determined by the following FOC (see (7)) for $k = 0$

$$-(1+\hat{\lambda}_k)q(1-q)v'(w_{NL}(\alpha_0^*(\hat{\lambda}_k)))E[X] + (1-q)E[v'(w_L(\alpha_0^*(\hat{\lambda}_k)))(X-(1+\hat{\lambda}_k)(1-q)E[X])] = 0,$$

and for $k > 0$

$$\begin{aligned}
& -(1+\hat{\lambda}_k)q(1-q)v'(w_{NL}(\alpha_k^*(\hat{\lambda}_k)))E[X](1+kg'(v(w_0)-v(w_{NL}(\alpha_k^*(\hat{\lambda}_k)))) \\
& + (1-q)E[v'(w_L(\alpha_k^*(\hat{\lambda}_k)))(X-(1+\hat{\lambda}_k)(1-q)E[X])(1+kg'(v(w_L^{\max})-v(w_L(\alpha_k^*(\hat{\lambda}_k)))))] \\
& = 0.
\end{aligned}$$

As $\alpha_k^*(\hat{\lambda}_k) = \alpha_0^*(\hat{\lambda}_k)$ those two conditions reduce to the following

$$\begin{aligned}
& -(1+\hat{\lambda}_k)q(1-q)v'(w_{NL}(\alpha_0^*(\hat{\lambda}_k)))E[X]g'(v(w_0)-v(w_{NL}(\alpha_0^*(\hat{\lambda}_k)))) \\
& + (1-q)E[v'(w_L(\alpha_0^*(\hat{\lambda}_k)))(X-(1+\hat{\lambda}_k)(1-q)E[X])g'(v(w_L^{\max})-v(w_L(\alpha_0^*(\hat{\lambda}_k))))] \\
& = 0.
\end{aligned}$$

As this condition is independent of k , $\hat{\lambda}_k = \hat{\lambda}$ for all $k > 0$ and thus $\alpha_k^*(\hat{\lambda}) = \alpha_0^*(\hat{\lambda})$ for all $k > 0$.

A.6 Proof of Lemma 6

In case of no loss, it is ex-post optimal to have bought no insurance, so $w_{NL}^{\max} = w_0$. In case of a realized loss of severity x , it is never ex-post optimal for the individual to buy insurance with a deductible level above x ,

since it would be more costly than not buying any insurance at all. In this case, the final level of wealth is $w_0 - x$. The optimal ex-post deductible level $\bar{D}(\lambda)$ below the realized loss x maximizes

$$w_L(D) = w_0 - (1 + \lambda)(1 - q)E[(X - D)^+] - D.$$

The first- and second-derivatives are

$$\begin{aligned} w'_L(D) &= (1 + \lambda)(1 - q)(1 - F(D)) - 1 \\ w''_L(D) &= -(1 + \lambda)(1 - q)f(D) < 0, \end{aligned}$$

where $f(\cdot)$ is the density function to $F(\cdot)$.

If $\lambda \leq \frac{q}{1-q}$, then $w'_L(0) \leq 0$ and concavity of $w_L(\cdot)$ implies $\bar{D}(\lambda) = 0$, i.e., full insurance is ex-post optimal. The ex-post optimal level of final wealth is then

$$w_L^{\max} = w_0 - \min(x, (1 + \lambda)(1 - q)E[X]). \quad (18)$$

If $\lambda > \frac{q}{1-q}$, then $w'_L(D) = 0$ if and only if

$$F(D) = \frac{\lambda(1 - q) - q}{(1 + \lambda)(1 - q)},$$

so

$$\bar{D}(\lambda) = F^{-1}\left(\frac{\lambda(1 - q) - q}{(1 + \lambda)(1 - q)}\right).$$

In this case, the ex-post optimal level of final wealth is

$$w_L^{\max} = \begin{cases} w_0 - x, & \text{if } x < P(\bar{D}(\lambda)) + \bar{D}(\lambda) \\ w_0 - P(\bar{D}(\lambda)) - \bar{D}(\lambda), & \text{if } x \geq P(\bar{D}(\lambda)) + \bar{D}(\lambda) \end{cases}. \quad (19)$$

A.7 Proof of Proposition 7

The individual solves the following optimization problem

$$\begin{aligned} \max_{D \in \mathbb{R}_+} E[u(w(D))] &= q(v(w_0 - P(D)) - kg(v(w_0) - v(w_0 - P(D)))) \\ &+ (1 - q) \int_0^D (v(w_0 - P(D) - x) - kg(v(w_L^{\max}) - v(w_0 - P(D) - x))) dF(x) \\ &+ (1 - q) \int_D^{w_0} (v(w_0 - P(D) - D) - kg(v(w_L^{\max}) - v(w_0 - P(D) - D))) dF(x). \end{aligned}$$

The first derivative yields

$$\begin{aligned}
\frac{dE[u(w(D))]}{dD} = & \\
& -qP'(D)v'(w_0 - P(D))(1 + kg'(v(w_0) - v(w_0 - P(D)))) \\
& - (1 - q)P'(D) \int_0^D v'(w_0 - P(D) - x)(1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - x))) dF(x) \\
& - (1 - q)(P'(D) + 1)v'(w_0 - P(D) - D) \int_D^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - D))) dF(x)
\end{aligned} \tag{20}$$

and the second derivative

$$\begin{aligned}
\frac{d^2E[u(w(D))]}{dD^2} = & \\
& -qP''(D)v'(w_0 - P(D))(1 + kg'(v(w_0) - v(w_0 - P(D)))) \\
& +qP'^2(D)v''(w_0 - P(D))(1 + kg'(v(w_0) - v(w_0 - P(D)))) \\
& -qP'^2(D)v'^2(w_0 - P(D))kg''(v(w_0) - v(w_0 - P(D))) \\
& - (1 - q)P''(D) \int_0^D v'(w_0 - P(D) - x)(1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - x))) dF(x) \\
& + (1 - q)P'^2(D) \int_0^D v''(w_0 - P(D) - x)(1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - x))) dF(x) \\
& - (1 - q)P'^2(D) \int_0^D v'^2(w_0 - P(D) - x)kg''(v(w_L^{\max}) - v(w_0 - P(D) - x)) dF(x) \\
& - (1 - q)P''(D)v'(w_0 - P(D) - D) \int_D^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - D))) dF(x) \\
& + (1 - q)(P'(D) + 1)^2v''(w_0 - P(D) - D) \int_D^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - D))) dF(x) \\
& - (1 - q)(P'(D) + 1)^2v'(w_0 - P(D) - D) \int_D^{w_0} kg''(v(w_L^{\max}) - v(w_0 - P(D) - D)) dF(x) \\
& + (1 - q)v'(w_0 - P(D) - D)(1 + kg'(v(w_L^{\max}(D)) - v(w_0 - P(D) - D)))
\end{aligned} \tag{21}$$

where, from equations (18) and (19), $w_L^{\max}(D) = w_0 - \min(D, (1 + \lambda)(1 - q)E[X])$ for $\lambda \leq \frac{q}{1-q}$ and $w_L^{\max}(D) = w_0 - \min\left(x, (1 + \lambda)(1 - q)E\left[(X - \bar{D}(\lambda))^+\right] + \bar{D}(\lambda)\right)$ for $\lambda > \frac{q}{1-q}$.

At $D = 0$, the first derivative is

$$\begin{aligned} & \frac{dE[u(W(D))]}{dD} \Big|_{D=0} \\ = & qP'(0)v'(w_0 - P(0))(1 + kg'(v(w_0) - v(w_0 - P(0)))) \\ & - (1 - q)(P'(0) + 1)v'(w_0 - P(0)) \int_0^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(0)))) dF(x). \end{aligned}$$

From (6) we get

$$\begin{aligned} & \frac{dE[u(w(D))]}{dD} \Big|_{D=0} \\ = & q(1 + \lambda)(1 - q)v'(w_0 - P(0))(1 + kg'(v(w_0) - v(w_0 - P(0)))) \\ & + (1 - q)(\lambda(1 - q) - q)v'(w_0 - P(0)) \int_0^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(0)))) dF(x). \end{aligned}$$

If $q > 0$ and $\lambda \geq \frac{q}{1-q}$, then

$$\frac{dE[u(w(D))]}{dD} \Big|_{D=0} > 0$$

so the individual can increase his expected utility by increasing the deductible, hence $D^* > 0$. If $q > 0$ and $\lambda < \frac{q}{1-q}$, then $w_L^{\max} < w_0$ and

$$\begin{aligned} & \frac{dE[u(w(D))]}{dD} \Big|_{D=0} \\ > & q(1 + \lambda)(1 - q)v'(w_0 - P(0))(1 + kg'(v(w_0) - v(w_0 - P(0)))) \\ & + (1 - q)(\lambda(1 - q) - q)v'(w_0 - P(0))(1 + kg'(v(w_0) - v(w_0 - P(0)))) \\ = & \lambda(1 - q)v'(w_0 - P(0))(1 + kg'(v(w_0) - v(w_0 - P(0)))) \\ \geq & 0. \end{aligned} \tag{22}$$

Combining these two results, we can conclude that if $q > 0$, then $D^* > 0$ for all $\lambda \geq 0$.

Now consider the case of $q = 0$. In this case,

$$\frac{dE[u(w(D))]}{dD} \Big|_{D=0} = \lambda v'(w_0 - P(0)) \int_0^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(0)))) dF(x) \tag{23}$$

which is positive if $\lambda > 0$. This implies that $D^* > 0$ if $\lambda > 0$.

If $q = 0$ and $\lambda = 0$ then

$$\frac{dE[u(w(D))]}{dD} \Big|_{D=0} = 0.$$

To determine whether $D = 0$ is optimal in this situation, we examine whether expected utility is a convex

or concave function in D at $D = 0$.⁴ For $q = 0$ and $\lambda = 0$, we derive from (6) and (20)

$$\begin{aligned} \frac{dE[u(w(D))]}{dD} = & \\ (1 - F(D)) \int_0^D v'(w_0 - P(D) - x) (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - x))) dF(x) & \\ - F(D) v'(w_0 - P(D) - D) \int_D^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - D))) dF(x) & \end{aligned} \quad (24)$$

where $w_L^{\max} = w_0 - \min(x, E[X])$ (see equation 19). Hence

$$\begin{aligned} \frac{d^2 E[u(w(D))]}{dD^2} & \\ = -f(D) \int_0^D v'(w_0 - P(D) - x) (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - x))) dF(x) & \\ + (1 - F(D))^2 \int_0^D v''(w_0 - P(D) - x) (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - x))) dF(x) & \\ - (1 - F(D))^2 k \int_0^D v'^2(w_0 - P(D) - x) g''(v(w_L^{\max}) - v(w_0 - P(D) - x)) dF(x) & \\ - f(D) v'(w_0 - P(D) - D) \int_D^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - D))) dF(x) & \\ + F^2(D) v''(w_0 - P(D) - D) \int_D^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D) - D))) dF(x) & \\ - F^2(D) v'^2(w_0 - P(D) - D) k \int_D^{w_0} g''(v(w_L^{\max}) - v(w_0 - P(D) - D)) dF(x) & \\ + f(D) v'(w_0 - P(D) - D) (1 + kg'(v(w_0 - \min(D, E[X])) - v(w_0 - P(D) - D))) & \end{aligned}$$

⁴Schlesinger (1981) emphasizes that expected utility can be a convex function in the deductible level over certain ranges.

At $D = 0$ we get

$$\begin{aligned}
& \frac{d^2 E[u(w(D))]}{dD^2} \Big|_{D=0} \\
&= -f(0) v'(w_0 - E[X]) \int_0^{w_0} (1 + kg'(v(w_0 - \min(x, E[X])) - v(w_0 - E[X]))) dF(x) \\
&\quad + f(0) v'(w_0 - E[X]) (1 + kg'(v(w_0) - v(w_0 - E[X]))) \\
&= f(0) v'(w_0 - E[X]) k \cdot \left(\begin{aligned} & g'(v(w_0) - v(w_0 - E[X])) \\ & - \int_0^{w_0} g'(v(w_0 - \min(x, E[X])) - v(w_0 - E[X])) dF(x) \end{aligned} \right) \\
&> 0
\end{aligned}$$

since $g'(v(w_0) - v(w_0 - E[X])) > g'(v(w_0 - \min(x, E[X])) - v(w_0 - E[X]))$ for all $x > 0$. For $q = 0$ and $\lambda = 0$ we thus have $\frac{dE[u(w(D))]}{dD} \Big|_{D=0} = 0$ and $\frac{d^2 E[u(w(D))]}{dD^2} \Big|_{D=0} > 0$. Convexity at $D = 0$ then implies that $D^* > 0$.

A.8 Proof of Proposition 8

This proposition is true if $\frac{dD_k^*(0)}{dk} > 0$. Applying the same approach shown in equation (12) we use the total differential to get,

$$\frac{dD_k^*(0)}{dk} = - \frac{\frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(0)}}{\frac{\partial^2 E[u(w(D))]}{\partial D^2} \Big|_{D=D_k^*(0)}}. \quad (25)$$

In the first step of this proof, we will show that $D_k^*(0)$ is an inner solution, (i.e. $0 < D_k^*(0) < w_0$), which implies local concavity (i.e., $\frac{\partial^2 E[u(w(D))]}{\partial D^2} \Big|_{D=D_k^*(0)} < 0$). In the second step, we will show that $\frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(0)} > 0$ which yields $\frac{dD_k^*(0)}{dk} > 0$. This will complete the proof.

We know from Proposition 7 that $D_k^*(0) > 0$. To show that $D_k^*(0) < w_0$, we will show that the RTEU function is convex at $D = w_0$ if $\lambda = 0$. Equation (21) implies

$$\begin{aligned}
& \frac{d^2 E[u(w(D))]}{dD^2} \Big|_{D=w_0} \\
&= -q(1-q) f(w_0) v'(w_0) (1 + kg'(0)) \\
&\quad - (1-q)^2 f(w_0) \int_0^{w_0} v'(w_0 - x) (1 + kg'(v(w_L^{\max}) - v(w_0 - x))) dF(x) \\
&\quad + (1-q) f(w_0) v'(0) (1 + kg'(v(w_0 - (1-q)E[X]) - v(0))).
\end{aligned}$$

Concavity of v implies $v'(0) > v'(w_0 - x)$ for all $0 \leq x < w_0$. Since $w_L^{\max} = w_0 - \min(x, (1-q)E[X])$ for all $x < (1-q)E[X]$,

$$g'(v(w_0 - (1-q)E[X]) - v(0)) > g'(v(w_L^{\max}) - v(w_0 - x)) = g'(0).$$

For all $x \geq (1-q)E[X]$,

$$g'(v(w_0 - (1-q)E[X]) - v(0)) > g'(v(w_L^{\max}) - v(w_0 - x)) = g'(v(w_0 - (1-q)E[X]) - v(w_0 - x)).$$

Therefore

$$\begin{aligned}
& \frac{d^2 E[u(w(D))]}{dD^2} \Big|_{D=w_0} \\
& > -q(1-q)f(w_0)v'(0)(1+kg'(v(w_0-(1-q)E[X])-v(0))) \\
& \quad - (1-q)^2 f(w_0)v'(0)(1+kg'(v(w_0-(1-q)E[X])-v(0))) \\
& \quad + (1-q)f(w_0)v'(0)(1+kg'(v(w_0-(1-q)E[X])-v(0))) \\
& = 0.
\end{aligned}$$

Expected utility is thus convex at $D = w_0$ for $\lambda = 0$ and hence $D_k^*(0) < w_0$. Since the maximum is an interior point, RTEU must be locally concave at $D_k^*(0)$. Therefore, equation (25) implies that

$$\text{sign} \left(\frac{dD_k^*(0)}{dk} \right) = \text{sign} \left(\frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(0)} \right).$$

Differentiating (20) with respect to k and setting $\lambda = 0$ yields

$$\begin{aligned}
& \frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(0)} = \\
& q(1-q)(1-F(D_k^*(0)))v'(w_0-P(D_k^*(0)))g'(v(w_0)-v(w_0-P(D_k^*(0)))) \\
& + (1-q)^2(1-F(D_k^*(0))) \\
& \times \int_0^{D_k^*(0)} v'(w_0-P(D_k^*(0))-x)g'(v(w_L^{\max})-v(w_0-P(D_k^*(0))-x))dF(x) \\
& + (1-q)((1-q)(1-F(D_k^*(0)))-1)v'(w_0-P(D_k^*(0))-D_k^*(0)) \\
& \times \int_{D_k^*(0)}^{w_0} g'(v(w_L^{\max})-v(w_0-P(D_k^*(0))-D_k^*(0)))dF(x).
\end{aligned}$$

The FOC $\frac{dE[u(w(D))]}{dD} \Big|_{D=D_k^*(0)} = 0$ implies that for $k > 0$,

$$\begin{aligned}
& \frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(0)} \\
& = -\frac{1}{k}(1-q) \cdot \left[\begin{aligned} & q(1-F(D_k^*(0)))v'(w_0-P(D_k^*(0))) \\ & + (1-q)(1-F(D_k^*(0))) \int_0^{D_k^*(0)} v'(w_0-P(D_k^*(0))-x)dF(x) \\ & + ((1-q)(1-F(D_k^*(0)))-1)v'(w_0-P(D_k^*(0))-D_k^*(0))(1-F(D_k^*(0))) \end{aligned} \right] \\
& > -\frac{1}{k}(1-q)v'(w_0-P(D_k^*(0))-D_k^*(0)) \cdot \left[\begin{aligned} & q(1-F(D_k^*(0))) \\ & + (1-q)(1-F(D_k^*(0)))F(D_k^*(0)) \\ & + ((1-q)(1-F(D_k^*(0)))-1)(1-F(D_k^*(0))) \end{aligned} \right] \\
& = 0.
\end{aligned}$$

Hence $\text{sign} \left(\frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(0)} \right) > 0$ and therefore $\frac{dD_k^*(0)}{dk} > 0$.

A.9 Proof of Proposition 9

This proof is in two parts. First we will show that for $k = 0$, the individual purchases no insurance at $\lambda = \bar{\lambda}$. Then we will show that at that same $\lambda = \bar{\lambda}$, the individual purchases some insurance when $k > 0$.

We begin by recalling from equation (6) that because $F(w_0) = 1$, $P'(w_0) = 0$. Therefore, from equation (20), $\frac{dE[u(w(D))]}{dD}|_{D=w_0} = 0$ for all $\lambda, q, k \geq 0$. We also know from the FOCs (22) and (23) that $\frac{dE[u(w(D))]}{dD}|_{D=0} > 0$ for all $\lambda > 0$, including $\lambda = \bar{\lambda}$. To show that $D_0^*(\bar{\lambda}) = w_0$, it is sufficient to prove that $\frac{dE[u(w(D))]}{dD} > 0$ for all $0 < D < w_0$. For $k = 0$ and $\lambda = \bar{\lambda}$, the first derivative (20) is

$$\begin{aligned}
& \frac{dE[u(w(D))]}{dD} \\
&= q(1-q)(1+\bar{\lambda})(1-F(D))v'(w_0-P(D)) \\
&\quad + (1-q)^2(1+\bar{\lambda})(1-F(D)) \int_0^D v'(w_0-P(D)-x)dF(x) \\
&\quad - (1-q)(1-(1-q)(1+\bar{\lambda})(1-F(D)))v'(w_0-P(D)-D)(1-F(D)) \\
&= (1-q)(1+\bar{\lambda})(1-F(D)) \left[qv'(w_0-P(D)) + (1-q) \int_0^D v'(w_0-P(D)-x)dF(x) \right. \\
&\quad \left. + (1-q)v'(w_0-P(D)-D)(1-F(D)) \right] \\
&\quad - (1-q)v'(w_0-P(D)-D)(1-F(D)) \\
&= (1-q)(1-F(D))v'(0) \cdot \frac{qv'(w_0-P(D)) + (1-q) \int_0^D v'(w_0-P(D)-x)dF(x) + (1-q)v'(w_0-P(D)-D)(1-F(D))}{qv'(w_0) + (1-q) \int_0^{w_0} v'(w_0-x)dF(x)} \\
&\quad - (1-q)v'(w_0-P(D)-D)(1-F(D)) \\
&= (1-q)v'(w_0-P(D)-D)(1-F(D)) \\
&\quad \times \left[\frac{\frac{v'(0)}{v'(w_0-P(D)-D)} \left(qv'(w_0-P(D)) + (1-q) \int_0^D v'(w_0-P(D)-x)dF(x) \right) + (1-q)v'(0)(1-F(D))}{qv'(w_0) + (1-q) \int_0^{w_0} v'(w_0-x)dF(x)} - 1 \right] \\
&> 0.
\end{aligned} \tag{26}$$

This inequality holds because $\frac{v'(0)}{v'(w_0-P(D)-D)} > 1$, $v'(w_0-P(D)) > v'(w_0)$, $v'(w_0-P(D)-x) > v'(w_0-x)$ for all x , and $v'(0)(1-F(D)) > \int_D^{w_0} v'(w_0-x)dF(x)$.

We have thus shown that at $\lambda = \bar{\lambda}$, expected utility is strictly increasing in D when $k = 0$. Therefore, $D_0^*(\bar{\lambda}) = w_0$ and it is optimal for the individual not to buy any insurance.

To show that a regret-averse individual ($k > 0$) buys some insurance at this loading factor ($D_k^*(\bar{\lambda}) < w_0$) we will prove that expected utility is convex at $D = w_0$. From equation (21), the second derivative at $D = w_0$

for $k > 0$ and $\lambda = \bar{\lambda}$ is

$$\begin{aligned}
& \frac{d^2 E[u(w(D))]}{dD^2} \Big|_{D=w_0} = \\
& -q(1+\bar{\lambda})(1-q)f(w_0)v'(w_0)(1+kg'(0)) \\
& - (1-q)(1+\bar{\lambda})(1-q)f(w_0) \int_0^{w_0} v'(w_0-x)(1+kg'(v(w_L^{\max})-v(w_0-x)))dF(x) \\
& + (1-q)f(w_0)v'(0)(1+kg'(v(w_L^{\max}(w_0))-v(0))) \\
& = - (1+\bar{\lambda})(1-q)f(w_0)k \cdot \left[qv'(w_0)g'(0) + (1-q) \int_0^{w_0} v'(w_0-x)g'(v(w_L^{\max})-v(w_0-x))dF(x) \right] \\
& + (1-q)f(w_0)v'(0)kg'(v(w_L^{\max}(w_0))-v(0)) \\
& = - (1-q)f(w_0)kv'(0) \cdot \frac{qv'(w_0)g'(0) + (1-q) \int_0^{w_0} v'(w_0-x)g'(v(w_L^{\max})-v(w_0-x))dF(x)}{qv'(w_0) + (1-q) \int_0^{w_0} v'(w_0-x)dF(x)} \\
& + (1-q)f(w_0)v'(0)kg'(v(w_L^{\max}(w_0))-v(0)) \\
& = (1-q)f(w_0)kv'(0)g'(v(w_L^{\max}(w_0))-v(0)) \cdot \left[1 - \frac{qv'(w_0)g'(0) + (1-q) \int_0^{w_0} v'(w_0-x)g'(v(w_L^{\max})-v(w_0-x))dF(x)}{g'(v(w_L^{\max}(w_0))-v(0)) \left(qv'(w_0) + (1-q) \int_0^{w_0} v'(w_0-x)dF(x) \right)} \right],
\end{aligned}$$

where

$$w_L^{\max}(w_0) = w_0 - (1+\bar{\lambda})(1-q)E[X] \text{ for } \bar{\lambda} \leq \frac{q}{1-q}$$

and

$$w_L^{\max}(w_0) = w_0 - (1+\bar{\lambda})(1-q)E[(X - \bar{D}(\bar{\lambda}))^+] + \bar{D}(\bar{\lambda}) \text{ for } \bar{\lambda} > \frac{q}{1-q}.$$

Hence for all $x < (1+\bar{\lambda})(1-q)E[X] (< (1+\bar{\lambda})(1-q)E[(X - \bar{D}(\bar{\lambda}))^+] + \bar{D}(\bar{\lambda})$ respectively) we have

$$g'(v(w_L^{\max}(w_0)) - v(0)) > g'(v(w_L^{\max}) - v(w_0 - x)) = g'(0)$$

and for all $x \geq (1+\bar{\lambda})(1-q)E[X] (\geq (1+\bar{\lambda})(1-q)E[(X - \bar{D}(\bar{\lambda}))^+] + \bar{D}(\bar{\lambda})$ respectively),

$$g'(v(w_L^{\max}(w_0)) - v(0)) > g'(v(w_L^{\max}) - v(w_0 - x)) = g'(v(w_L^{\max}(w_0)) - v(w_0 - x)).$$

Therefore,

$$\frac{qv'(w_0)g'(0) + (1-q) \int_0^{w_0} v'(w_0-x)g'(v(w_L^{\max})-v(w_0-x))dF(x)}{g'(v(w_L^{\max}(w_0)) - v(0)) \left(qv'(w_0) + (1-q) \int_0^{w_0} v'(w_0-x)dF(x) \right)} < 1,$$

and so $\frac{d^2 E[u(w(D))]}{dD^2} \Big|_{D=w_0} > 0$. Convexity at $D = w_0$ and the fact that $\frac{dE[u(w(D))]}{dD} \Big|_{D=w_0} = 0$ implies that the optimal deductible level is an interior point $D_k^*(\bar{\lambda}) < w_0$ for all $k > 0$.

Analogously to (25) we deduce

$$\frac{dD_k^*(\bar{\lambda})}{dk} = - \frac{\frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(\bar{\lambda})}}{\frac{\partial^2 E[u(w(D))]}{\partial D^2} \Big|_{D=D_k^*(\bar{\lambda})}}.$$

Expected utility is locally concave in the deducible level at $D = D_k^*(\bar{\lambda})$ as $D_k^*(\bar{\lambda})$ is an interior solution. $\frac{\partial^2 E[u(w(D))]}{\partial D^2} \Big|_{D=D_k^*(\bar{\lambda})} < 0$ therefore implies

$$\text{sign} \left(\frac{dD_k^*(\bar{\lambda})}{dk} \right) = \text{sign} \left(\frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(\bar{\lambda})} \right).$$

Differentiation (20) with respect to k and setting $\lambda = \bar{\lambda}$ yields

$$\begin{aligned} & \frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(\bar{\lambda})} = \\ & q(1-q)(1+\bar{\lambda})(1-F(D_k^*(\bar{\lambda})))v'(w_0 - P(D_k^*(\bar{\lambda})))g'(v(w_0) - v(w_0 - P(D_k^*(\bar{\lambda})))) \\ & + (1-q)^2(1+\bar{\lambda})(1-F(D_k^*(\bar{\lambda}))) \\ & \times \int_0^{D_k^*(\bar{\lambda})} v'(w_0 - P(D_k^*(\bar{\lambda})) - x)g'(v(w_L^{\max}) - v(w_0 - P(D_k^*(\bar{\lambda})) - x))dF(x) \\ & + (1-q)((1-q)(1+\bar{\lambda})(1-F(D_k^*(\bar{\lambda}))) - 1)v'(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda})) \\ & \times \int_{D_k^*(\bar{\lambda})}^{w_0} g'(v(w_L^{\max}) - v(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda})))dF(x). \end{aligned}$$

The FOC $\frac{\partial E[u(w(D))]}{\partial D} \Big|_{D=D_k^*(\bar{\lambda})} = 0$ implies for $k > 0$,

$$\begin{aligned}
& \frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(\bar{\lambda})} = \\
& = -\frac{1}{k} (1-q) (1+\bar{\lambda}) (1-F(D_k^*(\bar{\lambda}))) \cdot \left[\begin{aligned} & qv'(w_0 - P(D_k^*(\bar{\lambda}))) \\ & + (1-q) \int_0^{D_k^*(\bar{\lambda})} v'(w_0 - P(D_k^*(\bar{\lambda})) - x) dF(x) \\ & + v'(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda})) (1-F(D_k^*(\bar{\lambda}))) \end{aligned} \right] \\
& + \frac{1}{k} (1-q) v'(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda})) (1-F(D_k^*(\bar{\lambda}))) \\
& = -\frac{1}{k} (1-q) (1-F(D_k^*(\bar{\lambda}))) v'(0) \\
& \times \frac{qv'(w_0 - P(D_k^*(\bar{\lambda}))) + (1-q) \int_0^{D_k^*(\bar{\lambda})} v'(w_0 - P(D_k^*(\bar{\lambda})) - x) dF(x) + v'(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda})) (1-F(D_k^*(\bar{\lambda})))}{qv'(w_0) + (1-q) \int_0^{w_0} v'(w_0 - x) dF(x)} \\
& + \frac{1}{k} (1-q) v'(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda})) (1-F(D_k^*(\bar{\lambda}))) \\
& = -\frac{1}{k} (1-q) v'(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda})) (1-F(D_k^*(\bar{\lambda}))) \\
& \times \left[\frac{\frac{v'(0)}{v'(w_0 - P(D_k^*(\bar{\lambda})) - D_k^*(\bar{\lambda}))} \left(qv'(w_0 - P(D_k^*(\bar{\lambda}))) + (1-q) \int_0^{D_k^*(\bar{\lambda})} v'(w_0 - P(D_k^*(\bar{\lambda})) - x) dF(x) \right) + (1-q)v'(0)(1-F(D_k^*(\bar{\lambda})))}{qv'(w_0) + (1-q) \int_0^{w_0} v'(w_0 - x) dF(x)} - 1 \right] \\
& < 0,
\end{aligned}$$

where the last inequality follows from (26). Therefore $\text{sign} \left(\frac{\partial^2 E[u(w(D))]}{\partial D \partial k} \Big|_{D=D_k^*(\bar{\lambda})} \right) < 0$ and hence $\frac{dD_k^*(\bar{\lambda})}{d\bar{\lambda}} < 0$.

A.10 Proof of Proposition 10

For an arbitrarily fixed $k > 0$, $D_0^*(0) = 0 < D_k^*(0)$ and $w_0 = D_0^*(\bar{\lambda}) > D_k^*(\bar{\lambda})$. As all functions are sufficiently smooth, $D_k^*(\lambda)$ is a continuous function in λ and so there must exist a $\hat{\lambda}_k \in (0, \bar{\lambda})$ such that

$$D_k^*(\hat{\lambda}_k) = D_0^*(\hat{\lambda}_k).$$

As $D_k^*(\hat{\lambda}_k)$ is an inner solution, $\hat{\lambda}_k$ is determined by the following FOC (see equation 20) for $k = 0$

$$\begin{aligned}
& -qP'(D_0^*(\hat{\lambda}_k))v'(w_0 - P(D_0^*(\hat{\lambda}_k))) \\
& - (1-q)P'(D_0^*(\hat{\lambda}_k)) \int_0^{D_0^*(\hat{\lambda}_k)} v'(w_0 - P(D_0^*(\hat{\lambda}_k)) - x) dF(x) \\
& - (1-q)(P'(D_0^*(\hat{\lambda}_k)) + 1)v'(w_0 - P(D_0^*(\hat{\lambda}_k)) - D_0^*(\hat{\lambda}_k))(1-F(D_0^*(\hat{\lambda}_k))) \\
& = 0,
\end{aligned}$$

and for $k > 0$

$$\begin{aligned}
& -qP'(D_k^*(\hat{\lambda}_k))v'(w_0 - P(D_k^*(\hat{\lambda}_k)))(1 + kg'(v(w_0) - v(w_0 - P(D_k^*(\hat{\lambda}_k))))) \\
& - (1 - q)P'(D_k^*(\hat{\lambda}_k)) \int_0^{D_k^*(\hat{\lambda}_k)} v'(w_0 - P(D_k^*(\hat{\lambda}_k)) - x)(1 + kg'(v(w_L^{\max}) - v(w_0 - P(D_k^*(\hat{\lambda}_k)) - x)))dF(x) \\
& - (1 - q)(P'(D_k^*(\hat{\lambda}_k)) + 1)v'(w_0 - P(D_k^*(\hat{\lambda}_k)) - D_k^*(\hat{\lambda}_k)) \\
& \times \int_{D_k^*(\hat{\lambda}_k)}^{w_0} (1 + kg'(v(w_L^{\max}) - v(w_0 - P(D_k^*(\hat{\lambda}_k)) - D_k^*(\hat{\lambda}_k))))dF(x) \\
& = 0.
\end{aligned}$$

As $D_k^*(\hat{\lambda}_k) = D_0^*(\hat{\lambda}_k)$ those two conditions reduce to the following

$$\begin{aligned}
& -qP'(D_0^*(\hat{\lambda}_k))v'(w_0 - P(D_0^*(\hat{\lambda}_k)))g'(v(w_0) - v(w_0 - P(D_0^*(\hat{\lambda}_k)))) \\
& - (1 - q)P'(D_0^*(\hat{\lambda}_k)) \int_0^{D_0^*(\hat{\lambda}_k)} v'(w_0 - P(D_0^*(\hat{\lambda}_k)) - x)g'(v(w_L^{\max}) - v(w_0 - P(D_0^*(\hat{\lambda}_k)) - x))dF(x) \\
& - (1 - q)(P'(D_0^*(\hat{\lambda}_k)) + 1)v'(w_0 - P(D_0^*(\hat{\lambda}_k)) - D_0^*(\hat{\lambda}_k)) \\
& \times \int_{D_0^*(\hat{\lambda}_k)}^{w_0} g'(v(w_L^{\max}) - v(w_0 - P(D_0^*(\hat{\lambda}_k)) - D_0^*(\hat{\lambda}_k)))dF(x) \\
& = 0.
\end{aligned}$$

As this condition is independent of k , $\hat{\lambda}_k = \hat{\lambda}$ for all $k > 0$ and thus $D_k^*(\hat{\lambda}) = D_0^*(\hat{\lambda})$ for all $k \geq 0$.

References

- Baron, J. (2000). *Thinking and Deciding* (3rd ed.). Cambridge: Cambridge University Press.
- Baron, J. and J. Hershey (1988). Outcome bias in decision evaluation. *Journal of Personality and Social Psychology* 54(4), 569–579.
- Bell, D. E. (1982). Regret in decision making under uncertainty. *Operations Research* 30(5), 961–981.
- Bell, D. E. (1983). Risk premiums for decision regret. *Management Science* 29(10), 1156–1166.
- Briys, E. P. and H. Louberge (1985). On the theory of rational insurance purchasing: A note. *Journal of Finance* 40(2), 577–581.
- Connolly, T. and J. Reb (2003). Omission bias in vaccination decisions: Where's the "omission"? where's the "bias"? *working paper, University of Arizona*.
- Cummins, D., D. McGill, Winkelvoss, and R. Zelten (1978). *Consumer Attitudes Toward Auto and Home-owners Insurance*. Philadelphia, Pa.: Wharton School Department of Insurance, University of Pennsylvania.
- Doherty, N. A. and L. Eeckhoudt (1995). Optimal insurance without expected utility: The dual theory and the linearity of insurance contracts. *Journal of Risk and Uncertainty* 10(2), 157–179.

- Grace, M. F., R. W. Klein, P. R. Kleindorfer, and M. R. Murray (2003). *Catastrophe Insurance: Consumer Demand, Markets and Regulation*. Boston: Kluwer Academic Publishers.
- Holmstrom, B. (1979). Moral hazard and observability. *Bell Journal of Economics* 10(1), 74–91.
- Johnson, E. J., J. Hershey, J. Meszaros, and H. Kunreuther (1993). Framing, probability distortions and insurance decisions. *Journal of Risk and Uncertainty* 7, 35–51.
- Kahneman, D. and A. Tversky (1979). Prospect theory: An analysis of decision under risk. *Econometrica* 47(2), 263–291.
- Kunreuther, H. (2000). Insurance as a cornerstone for public-private sector partnerships. *Natural Hazards Review* 1, 126–136.
- Larrick, R. P. and T. L. Boles (1995). Avoiding regret in decisions with feedback: A negotiation example. *Organizational Behavior and Human Decision Processes* 63(1), 87–97.
- Loomes, G. (1988). Further evidence of the impact of regret and disappointment in choice under uncertainty. *Economica* 55(217), 47–62.
- Loomes, G., C. Starmer, and R. Sugden (1992). Are preferences monotonic - testing some predictions of regret theory. *Economica* 59(233), 17–33.
- Loomes, G. and R. Sugden (1982). Regret theory: An alternative theory of rational choice under uncertainty. *Economic Journal* 92(368), 805–824.
- Loomes, G. and R. Sugden (1987). Testing for regret and disappointment in choice under uncertainty. *Economic Journal* 97(Supplement), 118–129.
- Machina, M. J. (2000). Non-expected utility and the robustness of the classical insurance paradigm. In G. Dionne (Ed.), *Handbook of Insurance*, Huebner International Series on Risk, Insurance and Economic Security, pp. 37–96. Boston: Kluwer Academic Publishers.
- Mossin, J. (1968). Aspects of rational insurance pricing. *Journal of Political Economy* 76(4-1), 553–568.
- Pashigian, B. P., L. L. Schkade, and G. H. Menefee (1966). The selection of an optimal deductible for a given insurance policy. *Journal of Business* 39(1), 35–44.
- Quiggin, J. (1994). Regret theory with general choice sets. *Journal of Risk and Uncertainty* 8(2), 153–165.
- Razin, A. (1976). Rational insurance purchasing. *Journal of Finance* 31(1), 133–137.
- Rothschild, M. and J. E. Stiglitz (1976). Equilibrium in competitive insurance markets: An essay on the economics of imperfect information. *Quarterly Journal of Economics* 90(4), 629–649.
- Savage, L. J. (1954). *The Foundations of Statistics*. Wiley publications in statistics. New York: Wiley.
- Schlesinger, H. (1981). The optimal level of deductibility in insurance contracts. *Journal of Risk and Insurance* 48(3), 465–481.
- Schlesinger, H. (2000). The theory of insurance demand. In G. Dionne (Ed.), *Handbook of Insurance*, Huebner International Series on Risk, Insurance and Economic Security, pp. 131–151. Boston: Kluwer Academic Publishers.
- Slovic, P., B. Fischhoff, S. Lichtenstein, B. Corrigan., and B. Combs (1977). Preference for insuring against probable small losses: Insurance implications. *Journal of Risk and Insurance* 44(2), 237–258.
- Starmer, C. and R. Sugden (1993). Testing for juxtaposition and event-splitting effects. *Journal of Risk and Uncertainty* 6(3), 235–254.
- Sugden, R. (1993). An axiomatic foundation of regret. *Journal of Economic Theory* 60(1), 159–180.
- Yaari, M. E. (1987). The dual theory of choice under risk. *Econometrica* 55(1), 95–115.

- Zeelenberg, M. (1999). Anticipated regret, expected feedback and behavioral decision making. *Journal of Behavioral Decision-Making* 12(2), 93–106.
- Zeelenberg, M., W. W. van Dijk, A. S. R. Manstead, and J. van der Pligt (2000). On bad decisions and disconfirmed expectancies: The psychology of regret and disappointment. *Cognition and Emotion* 14(4), 521–541.