

Welfare consequences of the compound risks of index insurance

Glenn Harrison^{1,2} | Jimmy Martínez-Correa³ |
Karlijn Morsink⁴ | Jia Min Ng⁵ | J. Todd Swarthout⁶

¹Maurice R. Greenberg School of Risk Sciences and Center for the Economic Analysis of Risk, Robinson College of Business, Georgia State University, Atlanta, Georgia, USA

²Harrison is also affiliated with the School of Economics, University of Cape Town, Cape Town, South Africa

³Department of Economics, Copenhagen Business School, Frederiksberg, Denmark

⁴Department of Economics, Utrecht University. Morsink is also affiliated with the Development Economics Group, Wageningen University and the Center for the Economic Analysis of Risk, Georgia State University, Utrecht, Netherlands

⁵Center for the Economic Analysis of Risk, Robinson College of Business, Georgia State University, Atlanta, Georgia, USA

⁶Department of Economics, Andrew Young School of Policy Studies, Georgia State University, Atlanta, Georgia, USA

Correspondence

Karlijn Morsink, Department of Economics, Utrecht University. Morsink is also affiliated with the Development Economics Group, Wageningen University and the Center for the Economic Analysis of Risk, Georgia State University, Atlanta, GA, USA.
Email: k.morsink@uu.nl

Abstract

Index insurance is an attractive variant on the standard insurance contract that allows the determination of a loss event to be defined by one or more thresholds on an index that is positively correlated with actual losses. Index insurance also comes with a compound risk, basis risk. We examine how these compound risks for index insurance affect behavior towards the decision to purchase the product or not. Using incentivized experiments, we control the actuarial specifics of contracts offered, the information provided to decision makers, and our knowledge of the risk preferences of decision makers. We demonstrate that individuals that fail to process

This is an open access article under the terms of the [Creative Commons Attribution-NonCommercial-NoDerivs](https://creativecommons.org/licenses/by-nc-nd/4.0/) License, which permits use and distribution in any medium, provided the original work is properly cited, the use is non-commercial and no modifications or adaptations are made.

© 2025 The Author(s). *Journal of Risk and Insurance* published by Wiley Periodicals LLC on behalf of American Risk and Insurance Association.

compound risks, end up purchasing the contract more than other individuals, as well as more than they should, generating significant welfare losses. We also develop a natural information treatment that mitigates these welfare losses.

KEYWORDS

behavioral welfare economics, compound risk, consumer surplus, index insurance

JEL CLASSIFICATION

C90, D60, D81, D90

1 | INTRODUCTION

Index insurance is a variant on the standard insurance contract that allows the determination of a loss event to be defined by one or more thresholds on an index that is positively correlated with actual losses. In the case of a weather index, the measuring instrument could be a physical water measurement device or some statistical measure, and the thresholds could define a moderate, severe or extreme drought. Pre-defined payments would then be affected if the index threshold is triggered. One attraction of such insurance is that claims adjusters are not required to evaluate claim payment for every realized event. Traditional concerns with moral hazard and adverse selection are absent with index insurance products. Another attraction is that the index could entail a forecast event, allowing claims to be paid before the expected event.

The major complication of index insurance is that the product poses a compound risk for the agent considering whether to purchase it or not. One risk is that the index pays off, and the other risk is that there is, in fact, a loss. Considered jointly, these two risks imply an upside compound risk and downside compound risk. The effect of the downside compound risk is familiar from nonperformance risk for standard insurance contracts, even though it is part of the correct operation of an index insurance contract. We examine the manner in which these compound risks for index insurance affect behavior towards the decision to purchase the product or not. To effect control over actuarial parameters, the information provided to decision makers, and knowledge of the risk preferences of decision makers, we conduct laboratory experiments with subjects facing real financial consequences of their decisions. These laboratory experiments were designed to guide ongoing and future field experiments with index insurance.

A goal of our design is to highlight the role that compound risks play. We select the actuarial characteristics of 54 decisions to purchase the contract or not. These characteristics include the actual loss probability, the loss amount, the premium, and the probability that the index matches the actual loss of the decision-maker. We are also able to identify the extent to which each of our subjects violates the Reduction of Compound Lotteries (ROCL) axiom, simply by observing their choices over risky lottery pairs that define consistency with ROCL. We can also characterize the risk preferences of subjects more generally, using

econometric measures that assume Expected Utility Theory (EUT) or Rank-Dependent Utility (RDU).¹

Our experiments also feature two ways of characterizing the compound risks of index insurance for those considering whether to purchase it or not. In one treatment, we explain the manner in which the two risks of the index insurance contract operate, with numerical examples and simple logic. In the other treatment, we literally apply ROCL for the decision-maker, by additionally displaying the actuarially-equivalent (AE) simple lotteries that correspond to the compound lottery of the contract. In the former treatment, we seek to provide decision-makers with a way to evaluate compound risks for themselves; in the latter treatment, we directly provide decision-makers with the ROCL-consistent implications of evaluating compound risk, bypassing the need for the decision-maker to apply ROCL for themselves, and hence bypassing the effects of failures to apply ROCL consistently.

We evaluate behavior towards the purchase of index insurance from a descriptive and normative perspective. From our data, we can say whether the treatments differ in the extent to which agents purchase the index insurance contract, and the extent to which those choices are driven by their consistency with ROCL. These descriptive insights from our design are then put to use when we normatively evaluate the welfare consequences of the observed decisions of our subjects. The normative evaluation relies on our priors with respect to the risk preferences of subjects, beyond the data-based determination of their consistency with ROCL. We evaluate the certainty equivalent (CE) for each subject if they purchase the contract, and also evaluate the CE if they do not purchase the contract. The difference in these CE tells us whether the subject should purchase the contract or whether they should not purchase the contract, based on our priors over their risk preferences and, of course, the actuarial characteristics of the contract. The difference in CE also provides a natural metric for the *quantitative* welfare significance of correct and incorrect purchase decisions. Given the risk preferences of our subjects, our design has many contracts that should be purchased, as well as many that should not be purchased. The normative key is whether the observed purchase decisions match the decisions we should see, given the risk preferences and insurance contracts on offer.

The primary focus of our experimental design was the decision to purchase index insurance. The experimental design involves 202 student subjects who participated in two treatments: index insurance (II) with explanation of the logic of compound risk, and II with the additional display of the AE lottery. Each subject had to make 54 insurance choices, where they could decide to purchase insurance or not based on different loss amounts, loss probabilities, premium values, and index matching probabilities. Subjects were randomly assigned to one treatment. In the II treatment, 84 subjects made insurance decisions based on compound probabilities; in the AE treatment, 118 subjects received the same explanation of compound risks provided in II as well as information about AE probabilities. In the II treatment, subjects

¹The statement that insurance can be evaluated by the individual as a risk management tool by the change in the individual's subjective welfare from the decision to purchase the contract is, at one level, not controversial. As a simple matter of theory for standard models of risk preferences in economics, such as EUT, it is uncontroversial, as stressed by Harrison and Ng (2016) and Ericson and Sydnor (2017). However, the failure of the assumption that the individual does evaluate insurance in this manner, rather than the statement that the individual could have done so, is the basis of many descriptive analyses of insurance experiments. And the apparent failure of this assumption as a descriptive matter, along with the possibility that different models of risk preferences might be normatively unattractive, is the basis of many normative analyses of insurance experiments. Reviews of these analyses are provided by Harrison (2019), Harrison (2024) and Jaspersen et al. (2022).

were presented with information about personal event probabilities, index matching probabilities, and monetary outcomes. In the AE treatment, the decision process and realizations were the same as in the II treatment, but the display included “pie displays” that showed equivalent AE lotteries along with corresponding outcomes. Hence, treatment AE introduces an aid designed to assist individuals in processing compound risk.

A second goal of our experimental design was to identify directly the extent to which individuals violate ROCL. To do this, each subject made 10 lottery choices involving decisions between a simple lottery and a compound lottery, as well as 10 corresponding choices between the same simple lottery and a simple lottery that was actuarially equivalent to the compound lottery. We can then directly identify individuals who deviated from the ROCL axiom. For example, if someone chose a simple lottery over a compound lottery, but chose the AE simple lottery in the related pair, there would be a violation of ROCL. Our design allows us to say, based solely on observed choices of each individual, whether they exhibit 0, 1, 2, ..., 9, or 10 violations of the ROCL axiom out of a possible 10 violations.

A third goal of our experimental design was to generate priors for our evaluation of the welfare effects of observed index insurance purchases. To estimate the risk preferences of individuals for both simple risk and compound risk, each individual subject was presented 100 pairs of risky lotteries specifically designed to reveal risk preferences, assuming consistency with either EUT or RDU.

We show that individuals who struggle with processing compound risk, measured directly by violations of the ROCL axiom, excessively purchase insurance products with a risk of (perceived) contractual nonperformance, as compared to individuals who do not violate ROCL. However, when ROCL violators are given a decision aid that helps them process compound risk, they purchase less insurance, and their insurance purchase does not differ from individuals who don't violate ROCL.

We then show that ROCL violators generate welfare losses in the II treatment that does not include the decision aid. However, when the ROCL violators receive the decision aid, they do no worse in terms of welfare than individuals who do not violate ROCL. An important policy implication of our findings is that consumer welfare should increase if insurance supervisors and regulators require insurers to inform consumers, transparently and in a balanced manner, about the compounded probabilities of all potential states that insurance does and does not cover.

Out of the 10,908 decisions to purchase or not purchase insurance, we find that 59% are choices to purchase insurance. However, our structural welfare evaluation shows that, based on their risk preferences, subjects *should* only decide to purchase insurance in 50% of insurance decisions. A large share of welfare losses appears to be created by *excess purchase*. Subjects that violate ROCL consistently (22% of our subjects) display excess purchase that translate into welfare losses and, as hypothesized, our decision aid helps these subjects improve welfare.

In the treatment without the decision aid, subjects that violate ROCL chose to purchase insurance in 70% of the insurance decisions, as opposed to subjects that comply with ROCL who chose to purchase insurance in only 57% of the insurance choices, a difference that is statistically significant with a *p*-value less than 0.001. However, in the treatment with the decision aid, this difference in purchase decisions disappears, and both ROCL violators and ROCL compliers choose to purchase 59% of the insurance decisions.

This *excess purchase* of ROCL violators with respect to ROCL compliers in the II treatment translates into lower efficiency, which is improved by the decision aid in the AE treatment. Our structural welfare evaluation shows that the efficiency of choices by ROCL violators is

improved with the decision aid. In the II treatment, ROCL violators chose insurance “correctly” according to their risk preferences in 44% of the insurance decisions, and this proportion increased to 49% in the AE treatment, a difference that is statistically significant with a p -value of 0.01. Therefore, ROCL violators made insurance decisions that are more efficient, according to their own risk preferences, when they are offered the decision aid.

Our primary contribution is to the understanding of index insurance as a policy mechanism. Despite its potential, demand for index insurance *appears* low, and many interventions have focused on increasing index insurance take-up.² The contractual modifications inherent in index insurance imply, however, that the take-up of index insurance does not necessarily lead to an improvement in welfare for everyone. If that is the case, *low* demand may merely be a reflection of the fact that some consumers may be *better off* not purchasing the product.³ We contribute to this literature by demonstrating that, indeed, *excess* purchase of index insurance is an important driver of welfare losses. Many individuals who *do* take-up index insurance are actually worse-off in terms of welfare than they would have been without the insurance.⁴ In terms of implications for consumer protection and supervision of index insurance, the use of choice architectures that clearly communicate the inherent compound risk in the product can play a substantial role in enhancing welfare, particularly for those who struggle to process compound risk.

In Section 2, we present a conceptual framework with a model of the demand for insurance where information about compound probabilities that is presented to subjects differs. We also discuss the *general* logic of our experimental design, which is intended to operationalize this conceptual framework. In Section 3, we specify the details of the experimental design motivated by this framework. Section 4 presents descriptives of ROCL violations and the effects of our treatments on insurance purchase. We present our structural model of the welfare of insurance with compound risk, and our welfare results, in Section 5. Section 6 discusses the results, and in Section 7 we conclude.

2 | CONCEPTUAL FRAMEWORK

We first present a model of a demand function for index insurance where the information that is available to subjects about compound risk differs between the two treatment conditions. We then discuss our experimental design.

Assume an individual with a wealth endowment E and a personal loss probability p_L . If there is a loss, it has value L , leaving the individual with outcome $E - L$. If the individual does not purchase insurance, she faces the simple prospect $\{E - L, p_L; E, (1 - p_L)\}$ in the usual notation.

Now assume the individual is asked if she would like to purchase an index insurance contract or not. This contract might pay when the individual does not suffer a loss, and it might not pay when the individual does suffer a loss, depending on each case, on whether the index matches the actual loss of the individual. Assume that the probability that the index matches

²See Gaurav et al. (2011), Cole et al. (2013), Cole et al. (2014), Norton et al. (2014), Takahashi et al. (2016), Casaburi and Willis (2018), Belissa et al. (2019), and Ceballos and Robles (2020).

³See Clarke (2016).

⁴The same welfare-focused perspective for index insurance is provided by Carter and Chiu (2018) and Flatnes et al. (2018), using EUT assumptions on risk preferences.

the outcome of the individual is m . This matching probability implies a correlation ρ between the individual's loss and the index outcome equal to $1 - (2 \times (1 - m))$. Therefore, there are four possible outcomes, which we spell out in full to be explicit about the compound risks involved:

1. The individual pays the premium π , *experiences a loss of L* , and *the index differs* from the outcome of the individual. The insurance does not pay out and the individual ends up with $E - \pi - L$. The compound risk of the individual experiencing a loss and the index differing is $p_L \times (1 - m)$. This is the downside risk of the index insurance contract.
2. The individual pays the premium π , *experiences a loss of L* , and *the index matches* the outcome of the individual. Hence, the insurance pays out L , and the individual ends up with $E - \pi - L + L = E - \pi$. The compound risk of the individual experiencing a loss and the index matching is $p_L \times m$.
3. The individual pays the premium π , *experiences no loss*, and *the index matches* the outcome of the individual. Hence, the insurance does not pay out, and the individual ends up with $E - \pi$. The compound risk of the individual not experiencing a loss and the index matching is $(1 - p_L) \times m$.
4. The individual pays the premium π , *experiences no loss*, and *the index differs* from the outcome of the individual. Hence the insurance pays out L and the individual ends up with $E - \pi + L$. The compound risk of the individual not experiencing a loss and the index differing is $(1 - p_L) \times (1 - m)$. This is the upside risk of the index insurance contract.

Figure 1 displays these outcomes with specific parameter values from our experiments.

If the index insurance product is presented in its compound form, with the p_L and m probabilities separately identified, we have a demand function d_i for individual i that is a function of the parameters p_L, m, π, E, L :

$$d_i(p_L, m, \pi, E, L). \tag{1}$$

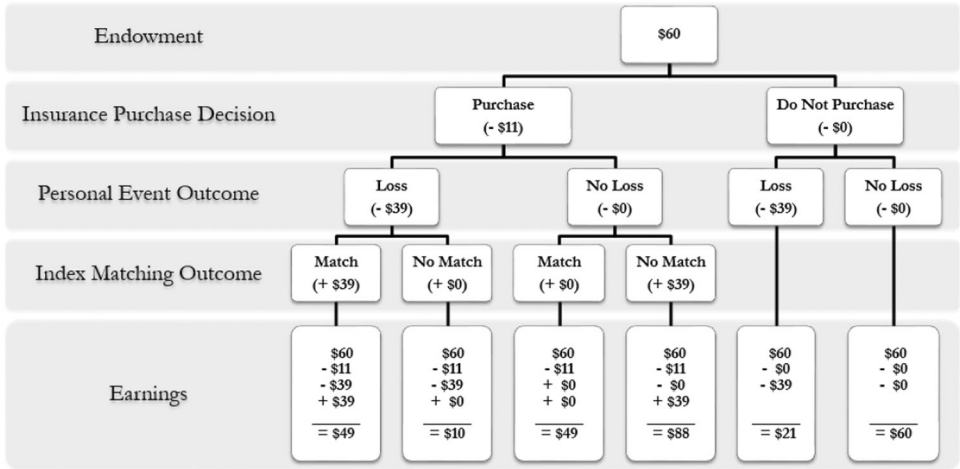


FIGURE 1 Decision tree for index insurance product. The decision tree and subsequent earnings for an index insurance product purchase assuming an endowment of \$60, a potential personal loss of \$39, and a premium of \$11 to purchase insurance to protect against that loss if the index matches.

This is the demand for index insurance that we would observe when the information presented to decision-makers is in compound form and the individual evaluates the insurance contract using EUT. The first two arguments of this demand function reflect the fact that the information given to the subject consists of the two layers of the compound risk in the index insurance: the probability of a personal loss and the probability of the index matching the personal loss.

In the AE treatment, we also provide the individual with a decision aid where we present the probabilities of the outcome states after multiplying the probability of the personal loss with the probability of the matching of the index. This implies that we provide information to the subject about the reduced form of the compound risk, the AE simple risk implied by the index insurance product. We then have a demand function D_i that is a function of the parameters for individual i :

$$D_i((p_L \times (1 - m)), (mp_L), ((1 - p_L) \times m), ((1 - p_L) \times (1 - m)), \pi, E, L). \quad (2)$$

This represents the demand for the index insurance contract when it is presented in reduced form, and the contract is again evaluated using EUT. Choices made using EUT imply ROCL-consistent choices.

Our experimental design is based on the hypothesis that individuals who behave consistently with the ROCL axiom will display the same insurance purchase behavior independently of the II and AE representation of the compound risk. Thus, when subjects do not violate the ROCL axiom, subjects behave as if $d_i = D_i$. We then hypothesize that the demand for index insurance of ROCL compliers will be the same across treatments.

However, if the individual violates ROCL, we hypothesize that the insurance purchase behavior in the representation with the decision aid is different from purchase behavior if the index insurance contract is also presented with probabilities in its compound form. It should be stressed that the AE treatment *augments* the information that the subject receives about probabilities, but does not change the underlying compound probability structure. The actual presence of compound risk is the same across treatments.

The general logic of our experimental design is guided by this simple framework. If a person satisfies ROCL, the demand for index insurance should not be affected by the decision aid because it only provides the information that the person should already know due to the ROCL assumption: the reduced form of the compound probabilities of each of the possible states of nature. Therefore, the index insurance demand for a person who satisfies ROCL should be the same with or without the decision aid.⁵

Now consider an individual that violates ROCL. A person that violates ROCL no longer finds equivalent risk presented in a compound form as presented in its reduced form. Therefore the prediction for such an individual is that index insurance demand across treatments could differ.

To quantify the quality of decisions, and potential gains from our decision aid, we use a measure of consumer welfare of insurance decisions defined in terms of the risk preferences of the individual. This allows us to measure whether an individual overinsures or underinsures, relative to the normatively optimal insurance purchase derived in terms of their risk preferences.

⁵In principle the provision of the decision aid could add some sort of “information overload,” and generate other effects on purchase decisions. Our subjects already had experience with simple lottery representations using exactly the form of the decision aid, so we do not believe that such overload would be plausible.

We calculate the CE of buying or not buying index insurance based on the actuarial parameters of a specific index insurance contract and the risk preferences of each individual. Consider the actual choices of one subject in our experiment, depicted in Table 1, displaying the 54 index insurance choices for this subject from the II treatment.⁶

Columns 2 through 6 of Table 1 display the actuarial parameters of the index insurance contract: the premium, probability of the loss, loss amount, loading, and matching probability, respectively. For the sake of the example, we assume that this subject's risk preferences can be described by the EUT model. We estimate the risk aversion utility parameter of 0.46 for this subject using the battery of risky choices that the subject made separately, and assuming the utility specification of $u(x) = x^{(1-r)}/(1-r)$; hence, our subject is modestly risk averse. This risk aversion parameter allows us to calculate the CE of buying insurance or not buying insurance, shown in columns 7 and 8. We then calculate the expected consumer surplus (CS), defined as the difference between the CE of buying insurance minus the CE of not buying insurance, which is shown in column 9. We have ordered the index insurance choices from the choice that gives the maximum expected CS to the one that gives the least expected CS; choices in the experiment were in random order for each subject. Column 10 shows the predicted choice according to the expected CS. If the CS is positive, the predicted choice is to buy insurance; if it is negative, the predicted choice is not to buy insurance. Column 11 shows the actual choices in the experiment for this subject.

We can also calculate an efficiency measure for each choice of each individual. Column 12 shows the *absolute value* of the expected CS in column 9 for subject 1. We show the absolute value because a person who is expected to buy insurance will have a CS gain equal to the difference of CE of buying insurance and not buying insurance, while if the person is expected not to buy insurance, the CS gain is the absolute value of this difference. Each and every choice in our insurance task provides the subject with an *opportunity* for a welfare gain. Column 13 shows the expected CS of actual choices. If this is positive, then the individual bought insurance when she was predicted to do so, as in choice #1. It can also be positive if the person did *not* buy insurance when she was predicted not to purchase insurance, as in choice #20. Alternatively, the CS can be negative if the person was predicted to buy insurance and did not, as in choice #3. This is the case that we define as *underinsurance*. Additionally, the CS can also be negative if she was predicted not to buy insurance, but still bought insurance, as in choice #30. This is the case that we define as *overinsurance*. This subject made 24 choices that generated positive consumer surplus, since they were the normatively "correct" choices according to the estimated risk preferences. Some of these normatively correct choices entailed buying insurance, and some entailed not buying insurance. However, this subject also made "wrong" choices, buying insurance when she should not have (overinsurance) and not buying insurance when she should have (underinsurance). These normatively correct or incorrect choices are flagged in column 14 as a welfare gain or loss.

Finally, an aggregate efficiency measure for this individual can also be calculated. We add up the CS measures of each actual choice in the experiment in column 13, and then divide that by the sum of the absolute value of predicted CS if all choices were made correctly. This gives a measure between -1 and 1 , which we normalize to be between 0% and 100% . For this individual, the measure of efficiency before normalization is equal to 0.006 , and equal to 50.3% after normalization.

⁶The choices are the same in the AE treatment, but just augmented with additional information.

TABLE 1 Ex Ante consumer surplus and predicted and actual for subject 1.

(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)
Choice #	Premium	Probability of loss	Loss amount	Loading	Matching probability	CE with insurance	CE without insurance	Consumer surplus (CS)	Predicted choices (1 = Insurance, 0 = No insurance)	Actual choices (1 = Insurance, 0 = No insurance)	Absolute value of CS	CS of actual choices	Welfare of choices (1 = Gain, 0 = Loss)
1	5	0.2	35	0.8	1	55	51.88	3.12	1	1	3.12	3.12	1
2	2.5	0.1	35	0.8	1	57.5	55.87	1.63	1	1	1.63	1.63	1
3	7	0.2	35	1	1	53	51.88	1.12	1	0	1.12	-1.12	0
4	6.25	0.2	35	0.8	0.9	48.2	47.28	0.92	1	1	0.92	0.92	1
5	1.5	0.1	20	0.8	1	58.5	57.83	0.67	1	0	0.67	-0.67	0
6	3.5	0.1	35	1	1	56.5	55.87	0.63	1	0	0.63	-0.63	0
7	2	0.1	20	1	1	58	57.83	0.17	1	0	0.17	-0.17	0
8	2.5	0.1	20	1.2	1	57.5	57.83	-0.33	0	1	0.33	-0.33	0
9	4.5	0.1	35	1.2	1	55.5	55.87	-0.37	0	1	0.37	-0.37	0
10	2.5	0.1	20	0.8	0.9	53.97	54.44	-0.47	0	0	0.47	0.47	1
11	9	0.2	35	1.2	1	51	51.88	-0.88	0	1	0.88	-0.88	0
12	4.5	0.1	35	0.8	0.9	48.62	49.56	-0.94	0	0	0.94	0.94	1
13	3.5	0.1	20	1	0.9	52.96	54.44	-1.48	0	0	1.48	1.48	1
14	7.75	0.2	35	0.8	0.8	41.33	42.88	-1.55	0	1	1.55	-1.55	0
15	3.75	0.1	20	0.8	0.8	49.28	51.14	-1.86	0	1	1.86	-1.86	0
16	9	0.2	35	1	0.9	45.34	47.28	-1.94	0	1	1.94	-1.94	0
17	4.75	0.1	20	1.2	0.9	51.7	54.44	-2.74	0	0	2.74	2.74	1
18	6.25	0.1	35	1	0.9	46.8	49.56	-2.76	0	1	2.76	-2.76	0
19	4.75	0.1	20	0.8	0.7	44.94	47.94	-3	0	0	3	3	1

(Continues)

TABLE 1 (Continued)

(1)	(2)	(3)	(4)	(5)	(6)	(7)	(8)	(9)	(10)	(11)	(12)	(13)	(14)
Choice #	Premium	Probability of loss	Loss amount	Loading	Matching probability	CE with insurance	CE without insurance	Consumer surplus (CS)	Predicted choices (1 = Insurance, 0 = No insurance)	Actual choices (1 = Insurance, 0 = No insurance)	Absolute value of CS	CS of actual choices	Welfare of choices (1 = Gain, 0 = Loss)
20	6.25	0.1	35	0.8	0.8	40.31	43.6	-3.29	0	0	3.29	3.29	1
21	5.25	0.1	20	1	0.8	47.76	51.14	-3.38	0	0	3.38	3.38	1
22	9.25	0.2	35	0.8	0.7	34.67	38.67	-4	0	0	4	4	1
23	6	0.1	20	0.8	0.6	40.45	44.83	-4.38	0	1	4.38	-4.38	0
24	11.75	0.2	35	1.2	0.9	42.44	47.28	-4.84	0	0	4.84	4.84	1
25	8.25	0.1	35	1.2	0.9	44.72	49.56	-4.84	0	1	4.84	-4.84	0
26	6.75	0.1	20	1.2	0.8	46.24	51.14	-4.9	0	0	4.9	4.9	1
27	6.75	0.1	20	1	0.7	42.91	47.94	-5.03	0	1	5.03	-5.03	0
28	11.25	0.2	35	1	0.8	37.54	42.88	-5.34	0	1	5.34	-5.34	0
29	7	0.1	20	0.8	0.5	36.34	41.83	-5.49	0	0	5.49	5.49	1
30	8.25	0.1	35	0.8	0.7	32.1	37.99	-5.89	0	1	5.89	-5.89	0
31	9	0.1	35	1	0.8	37.36	43.6	-6.24	0	1	6.24	-6.24	0
32	10.75	0.2	35	0.8	0.6	28.26	34.67	-6.41	0	0	6.41	6.41	1
33	8.5	0.1	20	1	0.6	37.91	44.83	-6.92	0	1	6.92	-6.92	0
34	8.75	0.1	20	1.2	0.7	40.87	47.94	-7.07	0	1	7.07	-7.07	0
35	10.25	0.1	35	0.8	0.6	24.34	32.74	-8.4	0	1	8.4	-8.4	0
36	13.25	0.2	35	1	0.7	30.18	38.67	-8.49	0	1	8.49	-8.49	0
37	10	0.1	20	1	0.5	33.29	41.83	-8.54	0	1	8.54	-8.54	0
38	12.25	0.2	35	0.8	0.5	22.16	30.87	-8.71	0	1	8.71	-8.71	0

TABLE 1 (Continued)

(1) Choice #	(2) Premium	(3) Probability of loss	(4) Loss amount	(5) Loading	(6) Matching probability	(7) CE with insurance	(8) CE without insurance	(9) Consumer surplus (CS)	Predicted choices (1 = Insurance, 0 = No insurance)		(11) Actual choices (1 = Insurance, 0 = No insurance)	(12) Absolute value of CS	(13) CS of actual choices	(14) Welfare of choices (1 = Gai- n, 0 = Loss)
39	14.5	0.2	35	1.2	0.8	33.93	42.88	-8.95	0	0	1	8.95	-8.95	0
40	11.75	0.1	35	1.2	0.8	34.36	43.6	-9.24	0	0	1	9.24	-9.24	0
41	11	0.1	20	1.2	0.6	35.36	44.83	-9.47	0	0	1	9.47	-9.47	0
42	12	0.1	35	1	0.7	27.96	37.99	-10.03	0	0	1	10.03	-10.03	0
43	12.25	0.1	35	0.8	0.5	17.15	27.86	-10.71	0	0	0	10.71	10.71	1
44	13	0.1	20	1.2	0.5	30.24	41.83	-11.59	0	0	0	11.59	11.59	1
45	15.5	0.2	35	1	0.6	22.76	34.67	-11.91	0	0	0	11.91	11.91	1
46	17.25	0.2	35	1.2	0.7	25.47	38.67	-13.2	0	0	0	13.2	13.2	1
47	14.75	0.1	35	1	0.6	19.31	32.74	-13.43	0	0	0	13.43	13.43	1
48	15.5	0.1	35	1.2	0.7	23.98	37.99	-14.01	0	0	1	14.01	-14.01	0
49	17.5	0.2	35	1	0.5	15.97	30.87	-14.9	0	0	1	14.9	-14.9	0
50	17.5	0.1	35	1	0.5	11.36	27.86	-16.5	0	0	0	16.5	16.5	1
51	20	0.2	35	1.2	0.6	17.1	34.67	-17.57	0	0	0	17.57	17.57	1
52	19	0.1	35	1.2	0.6	14.29	32.74	-18.45	0	0	0	18.45	18.45	1
53	22.75	0.2	35	1.2	0.5	8.86	30.87	-22.01	0	0	1	22.01	-22.01	0
54	22.75	0.1	35	1.2	0.5	5.05	27.86	-22.81	0	0	0	22.81	22.81	1
24														

Note: Values in columns 2, 4, 7, 8, 9, 12, and 13 refer to dollars.

The final measure of efficiency lies between 0% and 100%. The higher the measure, the more efficient the choices of an individual in the sense that the person is buying (or not buying) index insurance when the expected consumer surplus of buying (or not buying) is positive.

3 | EXPERIMENTAL DESIGN

To operationalize the conceptual framework, we use an experimental design consisting of tasks that allow us to identify the treatments and risk preference parameters. First, the subject completes a battery of 20 lottery choices that are 10 tests of the ROCL axiom for each subject, since each choice pair presents the risky choices in its compound risk form and separately in its simple risk form. We use these 10 tests to define who is a ROCL violator or ROCL complier. Next, the subject completes a battery of 80 risky lottery choices defined over simple risk. We use all 100 choices to identify the risk preferences that we use to calculate the CS measures of insurance choices.⁷ Finally, the subject completes 54 index insurance choices, where a random subset of subjects are shown these choices with the aid that reduces compound risk to one-stage simple risk. The other random subset of subjects are only shown these choices in the compound risk form. We then calculate the expected CS for each of the 54 insurance choices and compare these against the actual choices to calculate the efficiency measure of the individual over the 54 choices. We conduct experiments with 202 student subjects at Georgia State University.⁸

Subjects can choose to purchase the insurance or not, and at the end of the experiment one choice is randomly selected for payment. We provide an endowment E of \$60 for each choice. Loss amounts L are either \$20, \$30, \$35, or \$39, and loss probabilities p_L are either 0.1 or 0.2. Multiplicative premium loadings on the actuarially-fair premia could be 0.5, 0.8, 1, 1.08 and 1.2, resulting in a premium π . We include negative loadings because index insurance premia are often subsidized. Finally, the index matching probability m could be 1, 0.8, 0.6, 0.4, 0.2, and 0. For the loss probability of 0.1, we considered all variants of loss amounts, matching probabilities, and premium amounts. For the loss probability of 0.2, we considered all variations of matching probabilities and premium amounts for one of the loss amounts. Appendix B (available on request) contains all experimental instructions.

3.1 | Index insurance without the decision aid

In the II treatment, before subjects make their insurance purchase decisions, they receive instructions about the insurance. The index insurance decisions are presented as shown in Figure 2. The probability of a personal event and the probability of the index matching are presented separately to the subjects. The monetary outcomes are also presented, based on the outcomes of the personal event and the index matching. At the top of the screen, the initial endowment, the personal loss amount, the premium, and the claim payment in the event that the index is triggered, are also presented.

⁷All these 100 choices are risky lotteries that were randomized at the individual level. We also randomized the choices for the 54 insurance purchase decisions.

⁸In 2018, 102 subjects completed the experiments; in 2019, an additional 100 subjects were added.

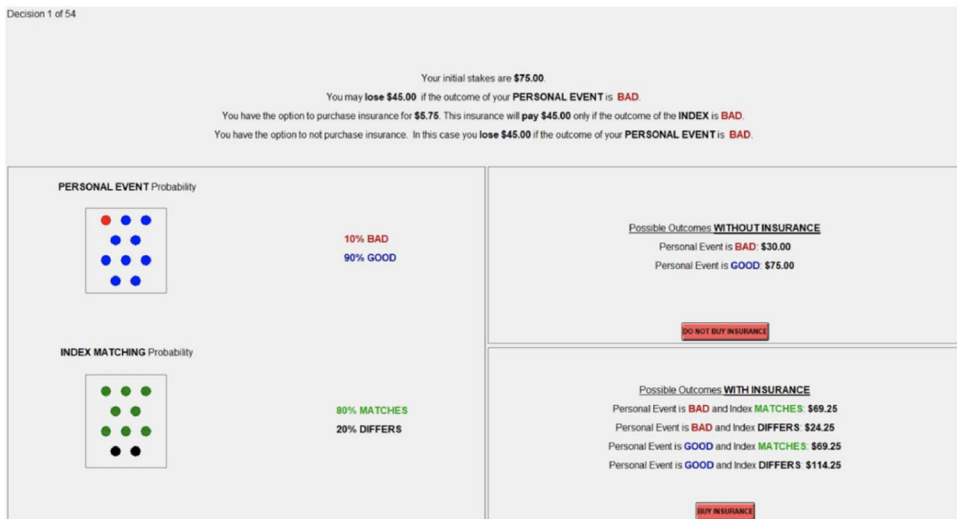


FIGURE 2 Example screen of index insurance without decision aid. [Color figure can be viewed at wileyonlinelibrary.com]

There are several important components of the logic of this task and the interface. The first is the use of the matching probability, m , between the personal loss and the index loss.⁹ The second component of the task, and the interface, is the use of distinct colors for the personal event (red and blue) and for the index matching (green and black). These colors are used to link the urns on the left to the payoffs on the right. To ensure the credibility of the random processes described, realizations were implemented by drawing appropriately colored chips from a “personal event” bag and an “index matching” bag. A third component of this task is the clear display of the two possible outcomes if the insurance is *not* purchased, and the four possible outcomes if the insurance is purchased. The four outcomes translate into only three distinct payoffs, but that redundancy is, we believe, valuable to help fully convey the operation of the product.¹⁰

3.2 | Index insurance with the decision aid

In the AE treatment task, the index insurance decisions, as well as the realizations based on drawing colored chips from a “personal event” bag and an “index matching” bag, are identical to the II task. The only difference is the display on the screens, adding “pie displays” showing

⁹One might have assumed that a simpler implementation would have been to specify a target correlation of the two, and randomly generate a personal and index realization consistent with that correlation. The logical difficulty is that one would need many such realizations for the subject to “experience” the intended sample correlation, and there would be some sample standard error around the experienced average correlation, raising additional issues of compound risk. Correlation is obviously a key actuarial parameter for this product. The method used here allows us to induce specific values for the correlation, and indeed to vary that from choice to choice.

¹⁰The only other experimental interface focused on index insurance that we are aware of was used in artefactual field experiments in Peru by Carter et al. (2008). Their design and interface were deliberately structured to mimic the field setting it was applied in, as a literacy treatment, whereas ours is deliberately structured to be more abstract, to allow evaluation of theoretical propositions. Each emphasis has a valid, complementary inferential role to play. Carter et al. (2008) do not report the results of the use of their literacy intervention.

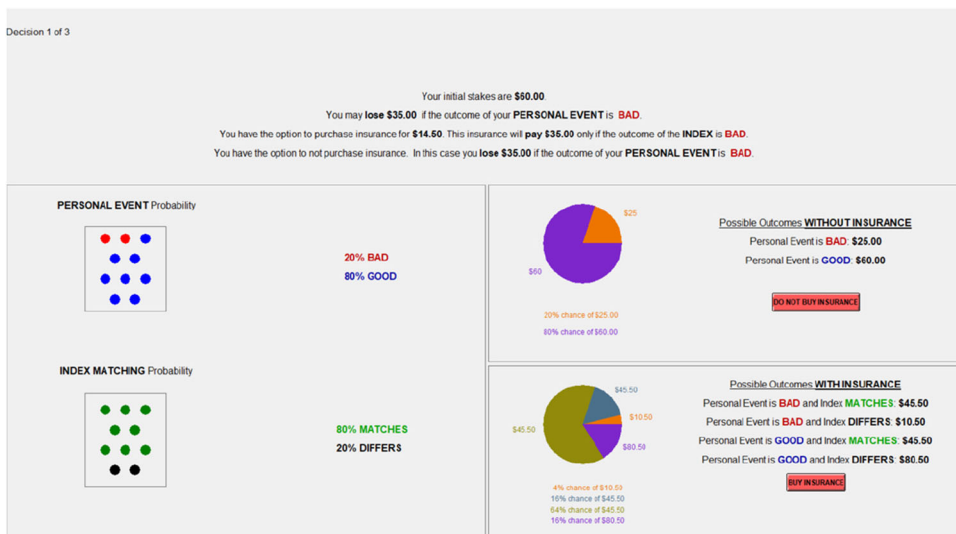


FIGURE 3 Example screen with actuarially-equivalent probabilities. [Color figure can be viewed at wileyonlinelibrary.com]

the equivalent AE lotteries implied, where the probabilities of the personal event and index matching are multiplied and presented in the pie display along with the corresponding outcome. Figure 3 shows an example screen. The instructions in the AE treatment were the same as for the II treatment, but complemented by extra information explaining the pie displays. The logic of the contract and underlying risk is explained in the same manner in the instructions for the II and AE treatments, so the natural context remains the same as in the II treatment.

All of the insurance choices came after a risky lottery task, and were presented in random order.¹¹ The average payoff per subject for the II and AE insurance task was \$57.88, with a standard deviation of \$12.49.

3.3 | Risky lottery choices

To be able to estimate the risk preferences of individuals over simple risk and compound risk, we use a task in which subjects are asked to make choices between 100 pairs of risky lotteries. In addition to its presentation of risk in a simple risk or compound risk form, we consider it to be attractive because it allows us to structurally estimate risk preferences for both EUT and RDU models, also taking into consideration that individuals make behavioral errors that we account for.

The 100 pairs of lotteries were designed to provide evidence of risk aversion as well as the tendency to make decisions consistently with EUT or RDU models. The battery is based on designs from Loomes and Sugden (1998) to test the Independence Axiom, designs from Harrison and Swarthout (2023) to evaluate Cumulative Prospect Theory (CPT)¹² models of risk

¹¹The insurance choices were programmed with the z-Tree software developed by Fischbacher (2007).

¹²Although we have risk lotteries that allow estimation of models of EUT, RDU, and CPT, we focus on the implications of estimating risk preferences using EUT and RDU. One reason is that there is little evidence of CPT behavior in this population, documented by Harrison and Swarthout (2023). Another reason is that we want to focus on different issues concerning the type of risk preference: whether one is an EUT or RDU decision-maker, and whether one obeys ROCL or not.

preferences, designs from Harrison et al. (2015) to test violations of the ROCL axiom, and a series of lotteries that are actuarially-equivalent versions of some of our index insurance choices.

3.4 | Measuring ROCL violations

To measure if individuals violate ROCL, Harrison et al. (2015) designed a revealed preference battery to non-parametrically test for violations of the ROCL axiom. Each subject was given 10 choices over a simple lottery and a compound lottery, as well as 10 corresponding choices over the same simple lottery and another simple lottery that was AE to the compound lottery. The two lotteries in these pairs were randomly assigned for presentation and were not presented contiguously.

This ROCL battery generates trade-offs in terms of foregone Expected Value (EV) that an individual who behaves consistently with ROCL would *not* care about. Consider lottery pair *rdon12* and lottery pair *rae12* in our battery.¹³ The Left and Right lottery in both *rdon12* and *rae12* are exactly the same, except that in *rdon12* the Right lottery is *presented* in compound form while in *rae12* it is *presented* in reduced form. In both *rdon12* and *rae12*, the Left lottery has an EV of \$43.75 and the Right lottery has an EV of \$52.50, with an EV difference of \$8.75. A person who satisfies ROCL would make consistent choices both in *rdon12* and *rae12*: if the person chooses the Left (Right) lottery in *rdon12*, then the person should choose Left (Right) lottery in *rae12*.

An individual who violates ROCL, however, will not have consistent choices when comparing *rdon12* and *rae12*. For example, say the individual chooses the Right lottery in the *rae12* lottery pair, but then chooses the Left lottery when presented with *rdon12*. This is a choice pattern inconsistent with ROCL. This behavior implies that the person is willing to forego \$8.75 in EV to avoid the Right (compound) lottery in favor of the Left (simple) lottery. This would be a violation of ROCL consistent with people displaying risk aversion towards compound risk that can be identified by observing how much the person is willing to forego in terms of EV to avoid a compound lottery that was previously preferred when presented in reduced form.

Our ROCL battery has 10 such lottery pair comparisons to identify ROCL violations, where we vary the difference in EV between the compound lottery and the simple lottery to provide trade-offs that allow us to identify the strength of attitudes towards compound risk.¹⁴ We count the number of the 10 pairs where each subject does *not* make ROCL-consistent choices, and use this as a measure of the degree to which each subject deviates from the ROCL axiom.¹⁵ This measure is agnostic about the reason for the choice inconsistency.¹⁶ For robustness, we also use a measure of ROCL violations that only identifies individuals who choose to consistently avoid compound risk as ROCL violators. These results are presented in Section 5.2.

¹³These are described numerically in Tables C.2 and C.4 in Appendix C (available on request).

¹⁴The EV differences of the compound lottery and simple lottery are: -\$17.50, \$0, \$1.25, \$2.50, \$7.50, \$8.25, \$20, and \$25.

¹⁵The compound lotteries are constructed by visually presenting two simple lotteries, but having some "double or nothing" option for one of them: Appendix C documents these 20 lottery pairs.

¹⁶There is some tendency for subjects in experiments to exhibit choice inconsistency even when faced with pairs of literally identical simple lotteries, and some of the ROCL violations could be due to that underlying inconsistency. Nonetheless, these are still violations of ROCL. Our econometric estimates of risk preferences account for this underlying choice inconsistency by means of a behavioral (Fechner) error parameter popularized by Hey and Orme (1994).

4 | DESCRIPTIVES AND INSURANCE PURCHASE

Figure 4 shows that our subjects violate ROCL between 0 and 9 times and that 78% makes 5 or fewer ROCL violations out of 10. We characterize individuals as ROCL-violators if they violate the ROCL consistent choice in more than half out of the 10 ROCL violation tests. Based on this categorization, 22% of the individuals in our sample are characterized as ROCL-violators.

Table 2 shows that our randomization to the II and AE treatments was successful in terms of observable characteristics. Columns 1 and 2 in Table 2 show the means and standard deviations of key covariates of the subjects in our sample by treatment assignment. The column “*t*-test difference” presents the difference in means between each of the treatments, and none of the differences are statistically significant (there are no asterisks beside any value). The F-test of joint significance is presented in the bottom row and is also not significant.

4.1 | Insurance purchase and ROCL violations

From our conceptual framework we have two hypotheses for insurance demand that we test. The first hypothesis is that the demand for insurance of ROCL violators may be affected by the

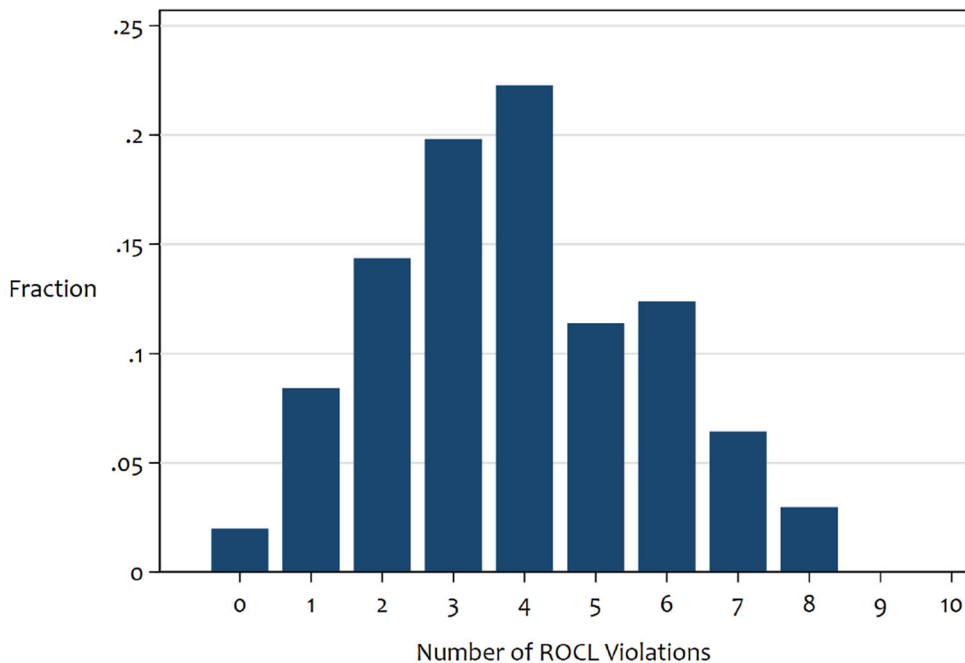


FIGURE 4 Histogram of the number of Reduction of Compound Lotteries (ROCL) violations per subject. Histogram of the number of ROCL violations per subject out of 10, based on a revealed preference measure of ROCL violations where each subject was given 10 lottery choices between a simple lottery and a compound lottery, as well as 10 corresponding lottery choices between the same simple lottery and a simple lottery that was actuarially-equivalent to that compound lottery. For each subject, we count the number of pairs out of the 10 where the subject does not make ROCL-consistent choices. [Color figure can be viewed at wileyonlinelibrary.com]

TABLE 2 Covariate means and balance.

Variable	AE Mean/(SD)	II Mean/(SD)	t-test difference
Female	0.627 (0.486)	0.560 (0.499)	−0.068
Age in years	28.805 (6.791)	27.167 (7.848)	−1.638
Black	0.669 (0.472)	0.726 (0.449)	0.057
Single	0.856 (0.353)	0.869 (0.339)	0.013
Household members	2.712 (1.904)	2.857 (1.889)	0.145
Business major	0.280 (0.451)	0.298 (0.460)	0.018
High GPA	0.585 (0.495)	0.631 (0.485)	0.046
Working part-time or full-time	0.593 (0.493)	0.619 (0.489)	0.026
Money spent per day	14.144 (10.599)	14.167 (10.689)	0.023
ROCL violations	3.898 (1.910)	3.798 (1.829)	−0.101
F-test of joint significance (F-stat)			0.499
Number of observations	118	84	202

Note: The value displayed for *t*-tests is the difference in the means across the groups. The value displayed for F-tests is the F-statistics. ***, **, and * indicate significance at the 1%, 5%, and 10% critical level. “Household members” refers to the number of household members in the household. “High GPA” represents a Grade Point Average score higher than 3.25 out of 4. “ROCL violations” refers to the number of ROCL violations out of 10.

treatments. The second hypothesis is that the demand for insurance of ROCL compliers is not affected by the treatments.

Table 3 displays the propensity to demand insurance by ROCL violators and ROCL compliers, and by treatment. We measure the propensity to demand insurance as the proportion of insurance choices that resulted in the purchase of index insurance. The first column allows us to test the first hypothesis about the demand of ROCL violators and the second column allows us to test the second hypothesis about ROCL compliers. As hypothesized in the conceptual framework, ROCL violators *are* affected by the framing of

TABLE 3 Percent of insurance choices of ROCL violators and ROCL compliers that resulted in insurance purchase.

	(1) ROCL Violators	(2) ROCL Compliers	(3) <i>p</i> -Value
II treatment	70%	57%	<0.001
AE treatment	59%	59%	0.78
<i>p</i> -Value	<0.001	0.03	

Note: The *p* values in the last column correspond to the *p*-value of a two-sided *t*-test that tests the hypothesis that proportions across ROCL violators and ROCL compliers are the same by treatment. The *p* values in the last row correspond to the *p*-value of a double-sided *t*-test that tests the hypothesis that proportions across treatments are the same for ROCL violators in column (1) and for ROCL compliers in column (2).

the underlying risk in the contract insurance. The first column of results in Table 3 shows that ROCL violators chose to buy insurance at a rate of 70% in the II treatment and at a rate of only 59% in the AE treatment, a difference of 11 percentage points. We can reject the hypothesis that these two proportions are different, with a *p*-value less than 0.001.

Additionally, we find that ROCL compliers are affected by the treatment, but this treatment effect is small compared to the treatment effect in ROCL violators. The second column of results in Table 3 shows that ROCL compliers chose to buy insurance at a rate of 57% in the II and at a rate of 59% in the AE treatment. A difference in means *t*-test shows that these two proportions are different, with a *p*-value of 0.03, but the economic significance of this difference (2 percentage points) is considerably smaller than the difference in purchase rate across treatments of ROCL violators (11 percentage points).

Table 3 also shows that in the II treatment, ROCL violators display higher purchase rates than ROCL compliers (70% vs. 57%). We can reject the hypothesis that these two proportions are the same, with a *p*-value less than 0.001. On the contrary, the purchase rate of ROCL violators is similar to ROCL compliers in the AE treatment, since the rate of insurance purchase is the same across treatments for these two types of subjects (59%). We cannot reject the hypothesis that these two proportions are the same, with a *p*-value of 0.78. We conclude that the decision aid in the AE treatment causes ROCL violators to display insurance purchase behavior similar to ROCL compliers.

We now analyze purchase behavior using statistical analysis with covariates. We estimate a panel logit model of purchase decision as a function of the indicator of ROCL violators, the treatment, the interaction of these two variables, and controlling for demographics and actuarial parameters, clustering errors at the level of the subject. Figure 5 displays the effect on the predicted probability of purchase of ROCL violators by treatment. ROCL violators are 16 percentage points more likely to purchase insurance in the II treatment, with a *p*-value of 0.01, while the probability of insurance purchase is not statistically different between ROCL violators and ROCL compliers in the AE treatment.

We conclude that the framing of the risk of index insurance affects the demand for insurance of people who violate ROCL. We present three pieces of evidence for this conclusion. First, the propensity to demand insurance of ROCL violators is higher in the II

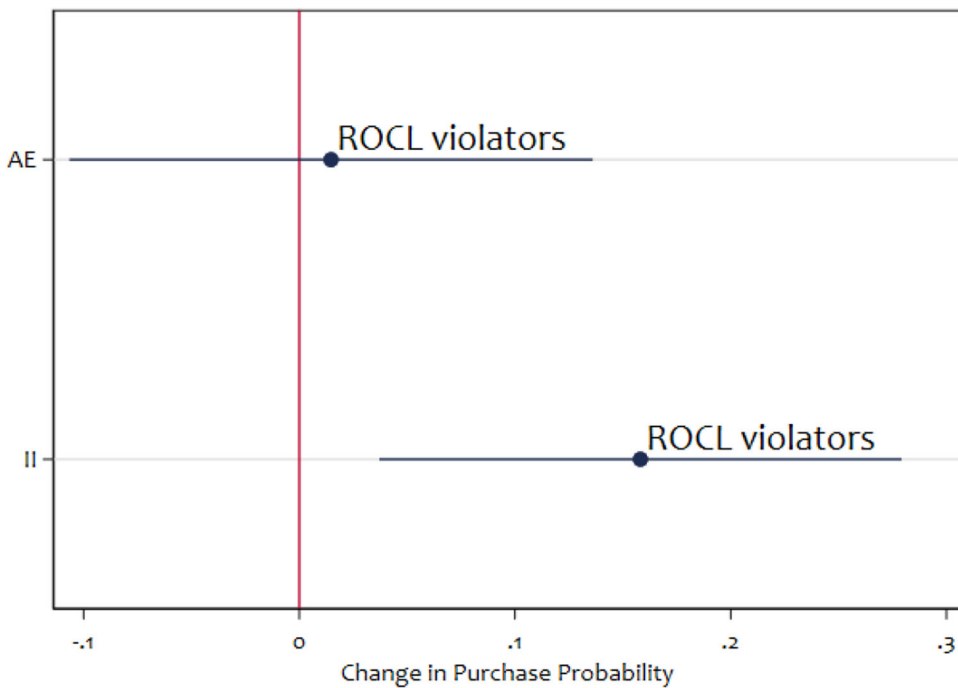


FIGURE 5 Effect of treatment on probability of insurance purchase by ROCL violators. The horizontal axis presents the predicted probability of purchase and 95% confidence intervals from panel logit regressions of the interactions of “ROCL violators” (vs. “non-ROCL violators”) and treatment on purchase. An individual is characterized as a ROCL violator if they violate ROCL in more than 5 out of the 10 tests for ROCL violations, and 22% of our subjects are classified as a ROCL violator. AE refers to the Actuarially Equivalent treatment, and II refers to the Index Insurance treatment. Regressions control for demographic and actuarial characteristics, and standard errors are clustered at the subject level. [Color figure can be viewed at [wileyonlinelibrary.com](https://onlinelibrary.wiley.com/terms-and-conditions)]

treatment than in the AE treatment, whereas the propensity to demand insurance of ROCL compliers is marginally affected by the treatment. Second, the propensity to demand insurance of ROCL violators is similar to the demand of ROCL compliers under the AE treatment, while the propensity to demand insurance of ROCL violators is higher in the II treatment. This evidence suggests that the AE decision aid helps ROCL violators with the processing of compound risk such that their insurance purchase behavior mimics the behavior of ROCL compliers. Third, statistical analysis that controls for demographics and actuarial parameters of contracts offered shows that subjects that violate ROCL are indeed the ones significantly affected by the treatments.

4.2 | Effects of actuarial parameters on insurance purchase

This effect of the AE treatment on insurance purchase of ROCL violators is sizeable compared to the effect of some of the traditional actuarial characteristics of the insurance contract. We run a logit model of purchase decision as a function of actuarial parameters: probability of loss, loss amount, loading, and matching probability. We focus on the last two relevant parameters. Figure 6 shows the marginal effects of an increase in the loading parameter, which we vary in

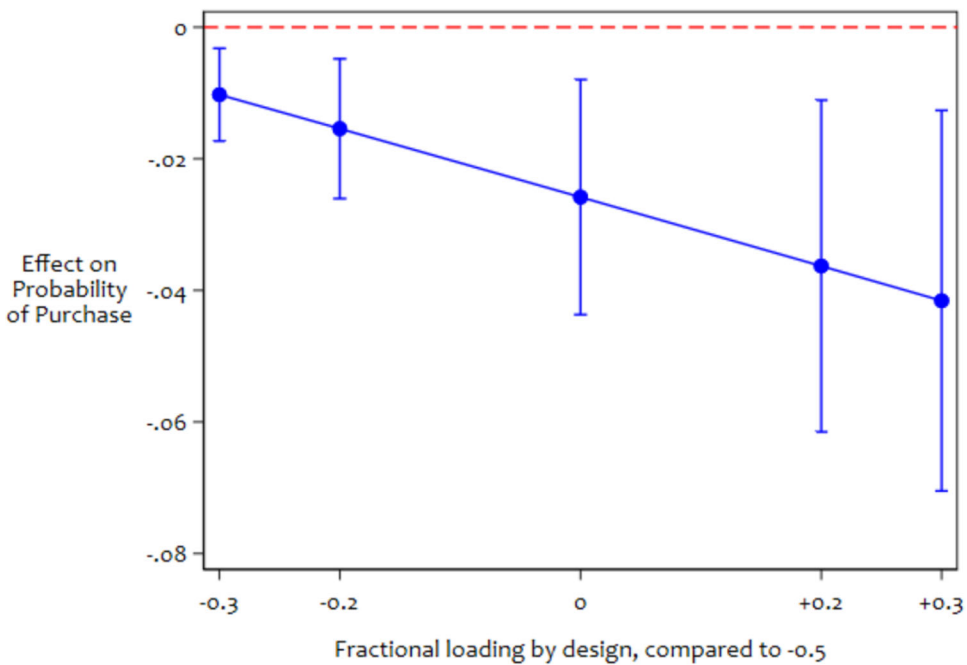


FIGURE 6 Effect of a higher loading parameter on purchase. The vertical axis presents the predicted marginal effect of insurance loading on the probability of purchase and 95% confidence intervals from a logit regression of the purchase decision on actuarial parameters and covariates. [Color figure can be viewed at wileyonlinelibrary.com]

the experiment between -0.5 and 0.3 .¹⁷ The x-axis of Figure 6 displays the loading level, and the y-axis displays the marginal effect of increasing the loading from -0.5 to a higher level. For instance, increasing the loading from -0.5 to 0.3 would decrease the probability of purchase by about 4 percentage points.

Another important actuarial parameter of an index insurance contract is the matching probability. Consider the marginal effect on the probability of purchasing the index insurance contract as the matching probability goes up. Figure 7 shows the marginal effects of an increase in the matching probability, which we vary from 0% to 100% in the experiment. The x-axis of Figure 7 displays the matching probability, and the y-axis displays the marginal effect of increasing the probability from 0% to a higher level. For instance, increasing the matching probability from 0% to 100% would *decrease* the probability of purchase by about 4.6 percentage points, which might appear to be counterintuitive.

This seemingly counterintuitive result provides an excellent illustration of the importance of recognizing heterogeneity of risk preferences when evaluating insurance contracts that entail compound risks, and hence more possible payoff outcomes. From the perspective of insurance as a risk management product, an increase in the matching probability entails a mean-preserving reduction in variance of outcomes, *ceteris paribus* the higher-order moments of the payoff distribution, and this should increase the likelihood of purchase. This is unambiguously

¹⁷This is the loading over actuarially-fair prices. A positive loading means that the premium is above the actuarially-fair premium, and a negative loading means that the premium is below the actuarially-fair premium.

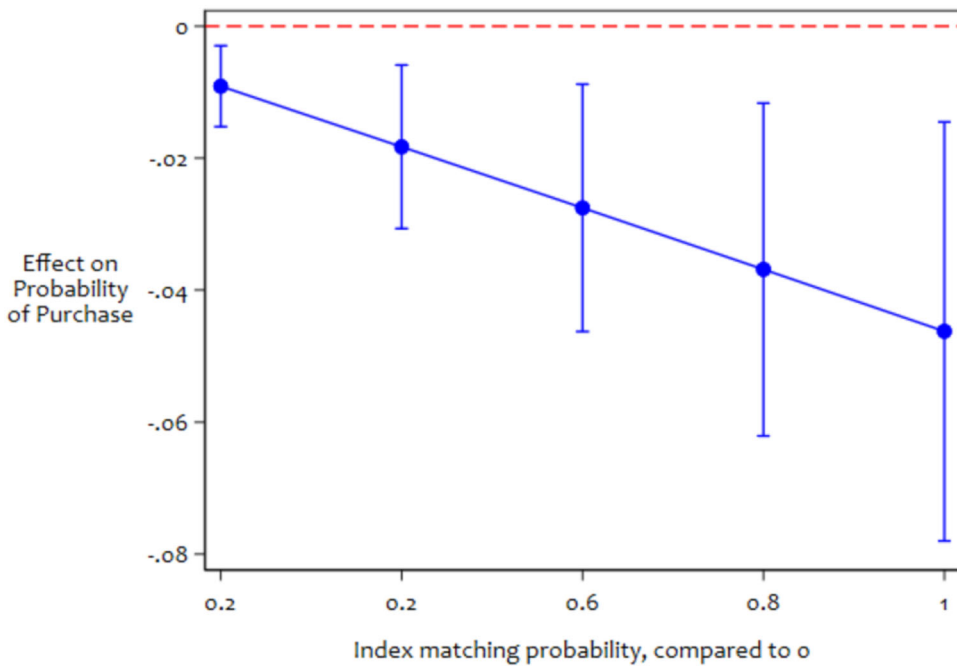


FIGURE 7 Effect of a higher matching probability parameter on purchase. The vertical axis presents the predicted marginal effect of the index matching probability on the probability of purchase and 95% confidence intervals from a logit regression of the purchase decision on actuarial parameters and covariates. [Color figure can be viewed at wileyonlinelibrary.com]

true for EUT maximizers with a concave utility function. However, this is not necessarily true for RDU maximizers.¹⁸

When probability weighting makes decision-makers globally optimistic or pessimistic for all probabilities, as is the case with the Power probability weighting function, the effects are relatively straightforward. Particularly important complications arise, however, when probability weighting allows for locally optimistic and locally pessimistic behavior towards different probabilities.¹⁹ Assume an Inverse-S probability weighting function $\omega(p) = p^\gamma / (p^\gamma + (1 - p)^\gamma)^{1/\gamma}$ for a RDU decision maker,²⁰ while decreasing the matching probability. This function exhibits inverse-S probability weighting (optimism for small p , and pessimism for large p) for $\gamma < 1$, and S-shaped probability weighting (pessimism for small p , and optimism for large p) for $\gamma > 1$. A smaller $\gamma < 1$ reflects an overweighting of the probabilities of extreme outcomes, while a larger $\gamma > 1$ reflects an underweighting of the probabilities of extreme outcomes.

Assume $\gamma > 1$. When no insurance is purchased, there are just *two* outcomes. The effect of S-shaped probability weighting, *ceteris paribus* any effect from $U'' < 0$, is to make the decision maker risk-averse with respect to the implied lottery when deciding not to purchase insurance.

¹⁸If one applies statistical tests that the probability weighting function for subjects involves no weighting, only 40% of our subjects would be classified as having risk preferences consistent with EUT using a significance level of 5% or less.

¹⁹Appendix A (available on request) explains and illustrates the logic of these complications.

²⁰This is just one of many popular specifications for the probability weighting function. Many people believe that “the” probability weighting function is the special case of this specification with $\gamma < 1$. This is simply not true, as stressed by Wilcox (2023).

The worst outcome, a loss, is given greater weight, and the best outcome, no loss, is given less weight.

However, and critically, when the matching probability is different from 1, the index insurance contract has *three* rank-ordered monetary outcomes: the “carrot” of no loss but a payout from the index, initial wealth minus the premium when the index matches the loss outcome, and the “stick” of a loss but no payout from the index. The carrot and stick here derive from the compound risks of index insurance. There is now a nuanced story when it comes to the implied lottery when deciding to purchase index insurance, since the purchase decision implies 3 outcomes rather than the 2 outcomes of the no-purchase decision. Assuming $\gamma > 1$, the intermediate payoff outcome, when the index matches the loss outcome, is given greater weight because of probability weighting. The worst outcome is given *slightly* smaller weight (almost imperceptibly smaller for $\gamma = 1.4$, for instance). However, the best outcome, the carrot of purchasing index insurance, is given *much* lower weight than the worst outcome. This is due to the asymmetric weighting of extremes in this instance.

Hence we have what can be called a “rotten carrot” effect from probability weighting: the lure of the good extreme outcome from basis risk is given less weight than it should have from the objective probabilities alone, and is also given less weight in a proportional sense than the curse of the bad extreme outcome from basis risk. Both effects of probability weighting serve to make the index insurance contract less attractive to somebody with these risk preferences, irrespective of the effects from $U'' < 0$. It is quite possible, even with low aversion to extremes from U'' , that the effect of probability weighting is to make the index insurance contract *less* attractive than facing the loss uninsured. And we know from Clarke (2016) that if there is also high enough aversion to extremes from U'' then the index contract could already be less attractive than being uninsured, even for an EUT decision maker.

What is the intuition for this seemingly counter-intuitive result? We expect less weight on the carrot outcome, *relatively* less weight on the carrot outcome than the stick outcome, and greater weight on the intermediate outcome. When the idiosyncratic loss probability is low, the complementary probability of no idiosyncratic loss is high, and this means that the carrot has a greater objective probability since it is one of the compound outcomes that depends on this complementary probability. Thus, the effect of probability weighting, to *proportionately reduce* this objective probability of the carrot outcome, has more (negative) impact on the overall evaluation of the lottery implied by purchasing the index insurance contract. When the loss probability gets large enough, this complementary probability gets smaller, hence the objective probability of the carrot outcome gets smaller, and the rotten carrot effect does not affect the overall evaluation of the lottery induced by purchasing the index insurance contract as much.

The bottom line is that one needs to know the level of risk aversion, defined by the risk premium under both EUT and RDU, to know if the index insurance product is attractive to purchase. But one also has to know the *type* of risk preference, to determine how much of the given risk premium derives from U'' and captures the benefits in “payoff variation reduction” of the insurance contract, and how much derives from the effect of probability weighting and models the behavioral evaluations of probabilities with which the index insurance contract pays out or not. Absent these two types of information, one cannot say *a priori* whether the take-up of a given index insurance product is welfare-enhancing for the individual.

These results provide another excellent illustration of the importance of recognizing heterogeneity of risk preferences when evaluating index insurance contracts. Enhancing

the matching probability is widely viewed as desirable from the perspective of making the index insurance contract more attractive, and considerable research is underway to employ remote-sensing technologies to complement local indices. However, such efforts could *reduce* the attractiveness of the index insurance product for certain RDU risk preferences. Specifically, S-shaped probability weighting can lead to this “rotten carrot” effect when the idiosyncratic loss probability is low, and “low” in this case includes the 10% or 20% loss probability levels in our experiments. The intuition that an improvement in the matching probability is surely an improvement in the quality of the index insurance contract runs squarely into the way in which compound lotteries are processed, and the way in which the probabilities of all payoffs are weighted under RDU. And the discussion of how RDU can lead to improved matching probabilities, making the index insurance product *less* attractive, alerts us to further pathways to the results in Figure 7.

5 | WELFARE

The objective of insurance is, *ex ante* any actual loss, to reduce the expected variability of consumption, by having the individual pay an insurance premium now in exchange for a claim payment later, in case the future state is realized where there is a loss. Therefore, we focus on an assessment of the expected welfare of buying insurance to an individual compared to the expected welfare of not buying insurance for the same individual, measured by the CS and efficiency measures defined in Section 2. For this evaluation, we estimate subjects preferences assuming RDU as it correctly embeds EUT if the individual does not display probability weighting.²¹

5.1 | Structural welfare evaluation

Our welfare analysis allows us to define when a subject purchases insurance when they should have not purchased, or to not buy insurance when they should have, according to their risk preferences. Table 4 shows the percentage of choices that resulted in insurance choices that were consistent with their individual risk preferences, hence “correct choices.” In the treatment without the decision aid, we found in Table 3 that subjects that violate ROCL chose to purchase insurance in 70% of the insurance decisions, whereas subjects that comply with ROCL chose to purchase insurance only in 57% of the insurance decisions. In the II treatment, Table 4 shows that ROCL violators should have purchased insurance in only 44% of the insurance decisions, while ROCL compliers should have purchased insurance in only 52% of the insurance decisions. However, in the AE treatment with the decision aid, this difference in “correct” purchase

²¹However, we do not follow the approach of Harrison and Ng (2016), classifying certain individuals as having risk preferences consistent with EUT or RDU. The statistical reason, explained by Monroe (2023), is that those subjects that are characterized as EUT by the test for “no probability weighting” still have standard errors around the probability weighting parameters, and potentially large ones. And, perhaps surprisingly, these standard errors can make a substantive difference in precisely the normative evaluations undertaken here. Hence, there is no formal need to differentiate EUT and RDU decision makers for these calculations, as noted by Gao et al. (2023), because EUT is nested within RDU, even if there is an important normative insight in knowing that there are these different types of risk preferences in the sample.

TABLE 4 Percent of choices of ROCL violators and ROCL compliers that resulted in insurance choices that were consistent with their risk preferences.

	(1) ROCL Violators	(2) ROCL Compliers	(3) <i>p</i> -Value
II treatment	44%	52%	<0.001
AE treatment	49%	49%	0.87
<i>p</i> -Value	0.01	0.04	

Note: The *p* values in the last column correspond to the *p*-value of a two-sided *t*-test that tests the hypothesis that proportions across ROCL violators and ROCL compliers are the same by treatment. The *p* values in the last row correspond to the *p*-value of a two-sided *t*-test that tests the hypothesis that proportions across treatments are the same for ROCL violators in column (1) and for ROCL compliers in column (2).

decisions disappears, and both ROCL violators and ROCL compliers choose, consistent with their risk preferences, to purchase insurance at the rate of 49%. This result shows that our decision aid effectively helps ROCL violators improve efficiency, such that the efficiency of their choices does not then differ from ROCL compliers.

Our aggregate measure of efficiency, estimated with the assumption of RDU preferences, tells a similar story. In the II treatment, the average efficiency across ROCL violators (39%) is lower than the average efficiency of ROCL compliers (50%), and this difference disappears with the AE treatment (46% average efficiency for ROCL violators and 47% average efficiency for ROCL compliers).

Using simple analyses of the choice pattern of insurance choices, as well as mean comparisons of “correct” choices and efficiency measures, we provide evidence that the *excess purchase* of ROCL violators with respect to ROCL compliers in the II treatment translates into lower efficiency that is improved by the decision aid in the AE treatment. We also repeat the statistical analysis that allows us to explore further the heterogeneity of the treatment effect across ROCL violators and ROCL compliers, while controlling for demographics and actuarial parameters of each contract offered to subjects.

Figure 8 shows that the significantly higher levels of purchase in the II frame by ROCL violators, documented in Table 3, translate into *significantly* lower levels of welfare for ROCL violators. Figure 8 displays the marginal effects estimated with a Beta regression that has the efficiency measure as the dependent variable and, as explanatory variables, the interaction of the dummies that identify ROCL violators with a dummy that identifies the AE and II treatment, as well as demographic and actuarial parameters. In the AE treatment there is no significant difference between ROCL violators and non-ROCL violators (*p*-value = 0.49), while there is a large and significant heterogeneous treatment effect in terms of welfare by whether or not individuals are ROCL violators. ROCL violators experience 21 percentage points lower Efficiency in the II treatment with respect to ROCL compliers, and this effect is significantly different from zero (*p*-value < 0.001).

We conclude that the higher propensity to purchase of ROCL violators in the II treatment translates into lower efficiency, and that the decision aid that we offer to individuals improves the efficiency of choices of ROCL violators.

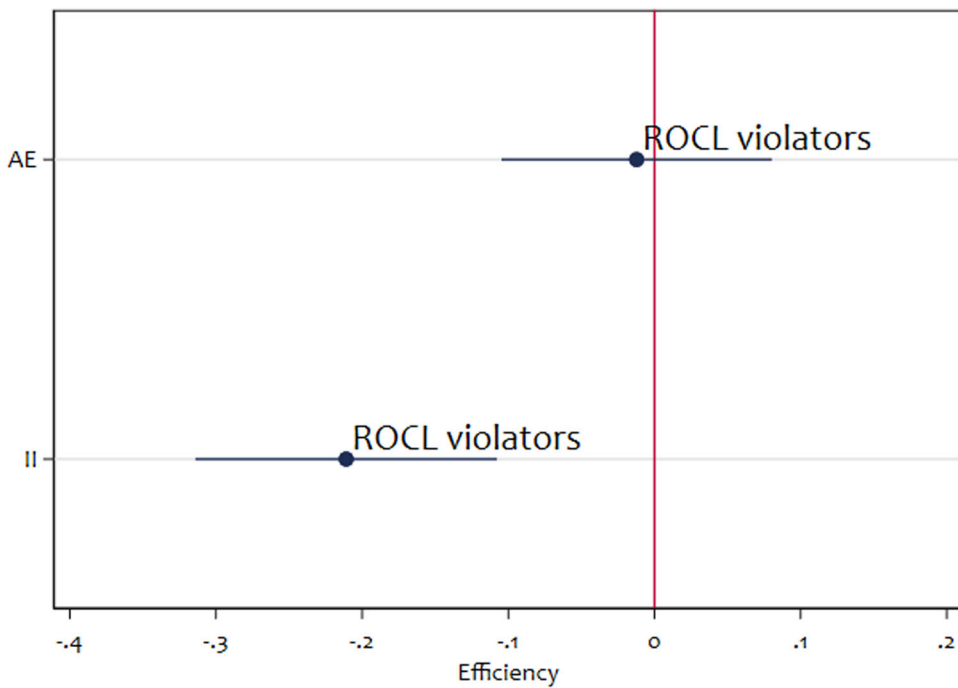


FIGURE 8 Welfare and ROCL violations. The horizontal axis presents the marginal effect on efficiency and 95% confidence intervals from Beta regressions calculated with RDU preferences of individuals who are classified as “ROCL violators” and “no ROCL violators.” An individual is characterized as a ROCL violator if they violate ROCL in more than 5 out of the 10 tests for ROCL violations. AE refers to the Actuarially Equivalent treatment, and II refers to the Index Insurance treatment. Regressions control for demographics and actuarial parameters, and standard errors are clustered at the individual level. [Color figure can be viewed at [wileyonlinelibrary.com](https://onlinelibrary.wiley.com/doi/10.1111/jori.12012)]

5.2 | Welfare evaluation without assuming ROCL

One conceptual limitation of our current methodology for calculating the expected welfare benefits from insurance is that we assume the subject calculates CS by using ROCL. We therefore consider a variant of the RDU model that does not assume ROCL.

The Recursive RDU model of Segal (1987) and Segal (1990) relaxes ROCL and allows probability weighting. It relaxes ROCL by assuming that the agent evaluates last-stage compound lotteries using RDU, replaces those lotteries with their RDU CE, then does the same for penultimate compound lotteries, and so on. Hence, the probabilities for all levels of compounding are “lost” in the sense that they are embedded in the CE, and the agent does not apply ROCL to them.²² This model can explain many of the EUT anomalies that RDU and other models were designed to explain, such as preference reversals (Segal, 1988), common ratio effects (Segal, 1990, §5), and the Ellsberg paradox (Segal, 1990). Moreover, this type of relaxation of ROCL affects the essence of “bundling or unbundling of risks,” which is at the heart of the higher-order risk preferences that characterize prudence and temperance

²²It is common in theoretical statements to use the CE, since then the first stage looks just like a simple lottery defined over deterministic outcomes.

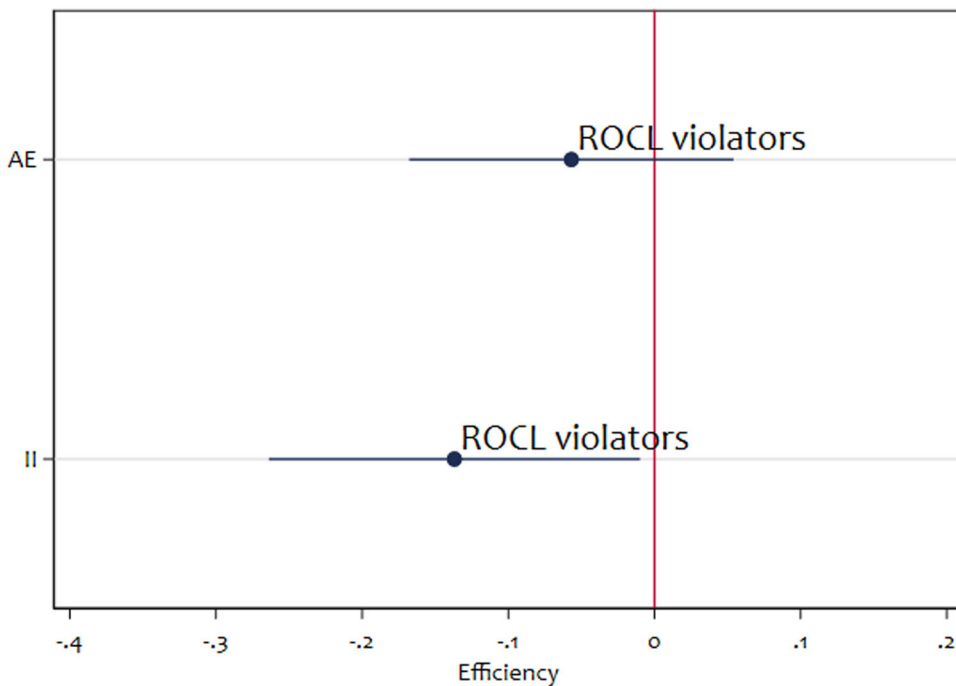


FIGURE 9 Welfare and ROCL violations. The horizontal axis presents the marginal effect on efficiency and 95% confidence intervals from Beta regressions calculated with *recursive* RDU preferences of individuals who are classified as “ROCL violators” and “no ROCL violators.” An individual is characterized as a ROCL violator if they violate ROCL in more than 5 out of the 10 tests for ROCL violations. AE refers to the Actuarially Equivalent treatment, and II refers to the Index Insurance treatment. Regressions control for demographics and actuarial parameters, and standard errors are clustered at the individual level. [Color figure can be viewed at wileyonlinelibrary.com]

(Eeckhoudt and Schlesinger, 2006). Quiggin (1993, p. 134) contrasts the standard, ROCL-consistent RDU model with the Recursive RDU model.

We undertake the same welfare analysis that assumed RDU risk preferences, but now assume the Recursive RDU risk preferences that allow us to model ROCL violations. Figure 9 shows that the calculation of expected welfare without assuming ROCL still leads to a similar conclusion as in Figure 8: the expected welfare gain is higher in the AE treatment than in the II treatment for the ROCL violators. Those ROCL violators have a 14 percentage points lower efficiency measure in the II treatment than subjects that are not classified as ROCL violators in the same treatment (p -value < 0.04). These results demonstrate again the importance of identifying heterogeneity in risk preferences when evaluating welfare gains and losses of insurance purchases.

6 | DISCUSSION

One limitation of our study is that it is not a field application of decision aids in the index-insurance context. In the field the probabilities of loss and the compound probabilities of index matching may be difficult to estimate and convey to people. It is then unclear whether the same issues with ROCL, treated as an objective risk, map as neatly into environments with subjective beliefs, and some uncertainty or even ambiguity.

An obvious response is that understanding these ideas in a controlled laboratory environment is a sensible first step toward understanding them in field settings. Indeed, index insurance in the field is designed based on historical realizations of an index, leading to thresholds that are clearly defined in a contract. Individuals who have lived through that history will have subjective beliefs about losses that jointly inform their beliefs about the probability of the index matching, although there is no presumption that those subjective beliefs statistically match the historical record. The ROCL axiom implies that one has to have clarity about the tree of possibilities or possible states of nature before one can actually calculate any probability to each branch of this tree. In the field, these could be elicited, communicated, and used in decision aids.

We also agree that many factors other than compound risk will determine index insurance purchase in a field setting, although there is empirical evidence from the field that shows that perceptions of basis risk matter for purchase behavior (e.g., Mobarak and Rosenzweig, 2013). And these additional factors make it even more important that we start the exercise of evaluating compound risks with the simplest, two-stage compound risks. From our experience in the field we would add subjective perceptions of personal loss probability, subjective perceptions of the loss amount when losses are defined in terms of physical outcomes (e.g., loss of goats or cows), and the perennial subjective risks of non-performance of the contract after the fact.²³ To the extent that *each* of these risks involves subjective beliefs, and those beliefs in turn take the form of subjective belief distributions, compound risk is also involved in their processing: see Harrison (2011, §4). We also stress again that the compound risks that are inherent to a fully-functioning index insurance contract should not be confused with contract nonperformance. Those compound risks, captured in our matching probability, arise when the contract is “performing” perfectly as a statistical measure of personal probability. Nonperformance risks, which arise from things such as fraud, the exploitation of legal “fine print,” or inadequate reserving, are certainly additional compound risks to any insurance contract, not just index insurance contracts. So these additional compound risks all entail multi-stage compounding, which, to our knowledge, have not been studied in the laboratory with incentivized choices.

Spelling out the tree of possibilities is an important aid to understand the payoff structure of all insurance. Providing probabilities to each branch of the tree is an added bonus, but we conjecture that the individual benefits greatly from just having clarity about the *possible* states of nature. There is precedence to this mechanism in the insurance industry itself. Reinsurers have to deal with what they call “cascade of events.” For example, an earthquake can be a trigger event that generates a tsunami that can kill people and destroy property, but it can also flood a nuclear power station, which in turn can create a nuclear meltdown. This cascade of events is “just” compounded risk, and reinsurers have mathematical models that give them probability distributions for each of these layers of risk. Of course, it is important to understand the probabilities of each branch of this cascade of events, but at some point these probabilities

²³Although not so important for index insurance, an additional factor for traditional insurance arises when the contract has exclusions and conditions for coverage. It is common for consumers to think that their losses are fully covered, or unconditionally covered, when in fact exclusions and underwriting conditions place limits to that coverage. An important example during COVID-19 was that insurance for “loss of business” did not include losses caused by pandemics, due to previous experience that underwriters had with the risks posed by Avian flu pandemics. Although coverage for pandemics is often provided through a rider, which should have flagged that it was not covered otherwise, many small businesses during COVID-19 expected to be covered for business losses. Similarly, many condo owners facing hurricane-generated closures of their condo building fail to realize that losses from rental income are not covered by condo homeowners insurance unless specifically paid for as a rider.

become so small that it becomes inefficient to speculate about probabilities. Instead, reinsurers use the mathematical models to generate scenarios that they do not know the probabilities of, and check if the reinsurance company can survive those scenarios with planned reserves. If this is not the case, the cautious thing to do for a reinsurer is then to include an exclusion in the reinsurance contract for such unlikely scenarios. Hence, thinking about the tree of possibilities itself, even without thinking about probabilities, can be informative for reinsurers. In our experiment, the aid we offer to people is spelling out to them the trees of possibilities, being clear about the scenarios the index insurance pays or does not pay out, and their respective probabilities. Our normative view is that the principle of the aid, spelling out all possible states of nature and its probabilities, helps people, just as with reinsurers, understand the risk they are exposed to, and the circumstances under which the insurance does or does not pay out. This information should, in turn, translate into better decisions for the individual.

How robust are our results to various definitions of a subject being a ROCL violator? We have considered the use of ROCL violations as a variable that takes on all 11 possible values, as distinct from our binary definition of a ROCL violator. To examine this issue, we created 6 binary definitions in which the subject had 2 or more, 3 or more, up to 7 or more violations; we do not observe more than 8 violations (Figure 4). We then re-estimated the logit model for purchase and the Beta regression for efficiency, using each of these definitions. Roughly 65% of the subjects exhibit 3, 4, 5, or 6 violations, suggesting that these would be the best candidates for examination.

Examining the model of purchase probability, we find that the II treatment effect estimate is positive across all the definitions of ROCL violation. However, it is only statistically significant with a p -value below 0.05 for the definitions using 5 or more and 6 or more violations. Examining the model of efficiency assuming RDU, we see a similar story. With the classic RDU model of risk preferences, the effect of II on the efficiency of ROCL violators is negative and statistically significant, with p -values below 0.05 for the definition using 5 or more and 6 or more violations. With the recursive RDU model of risk preferences, we have the same result as for the classic RDU model, although we also find the same significant negative effect for the definition using 4 or more violations.

Although our results are robust to different definitions of ROCL that we have sufficient data to statistically evaluate, we learn that for future (lab or field) experiments, we need to find better ways to identify the extent of ROCL violation at the level of the individual. Simply increasing the size of the ROCL battery that each subject completes would be one way to accomplish that, but likely costly in the field, where time required for such experimental tasks is usually constrained. Another alternative, demonstrated by Gao et al. (2023, §3.1), would be to keep the size of the ROCL battery for each subject fixed at, say, 10 pairs (hence 20 choices), but to select those 10 pairs for each subject at random from a wider battery, and then use a Bayesian Hierarchical Model to draw more reliable inferences about ROCL-consistency at the individual level. This alternative would trade off the statistical benefit of having ROCL-consistency determined by data, the choice patterns for each ROCL pair. But it would allow a wider range of tests of ROCL (e.g., three-stage or four-stage compounding, as well as two-stage compounding).

7 | CONCLUSION

We show that individuals who struggle with processing compound risk, measured directly by violations of the ROCL axiom, excessively purchase index insurance products, leading to expected welfare losses. However, when these ROCL violators are given a decision aid that helps

them process compound risk, they purchase less insurance and do no worse in terms of expected welfare than individuals who do not violate ROCL. The aid that we use is simple: we multiply out probabilities of the compound layers of risk, and explicitly inform subjects about the eventual likelihood of each outcome. An important policy implication of our findings is that average consumer welfare should increase if insurance supervisors and regulators require insurers to inform consumers about the compounded probabilities of all potential states that an index insurance product does and does not cover. Our results and these policy implications should apply more broadly to any contract with compound risk, such as pensions and warranties.

The general point is that we should focus on the literacy required to understand index insurance contracts, much as we should for insurance contracts in general. Index insurance serves to explicitly “force” attention on compound risk, but it is present in all insurance contracts in the field. Moreover, there are rigorous tools for documenting the extent of insurance literacy, as demonstrated by Harrison et al. (2022) for index insurance contracts. Once these literacy measures are in place, it is natural to document how decision aids, such as the aid examined here, change them, and ideally, improve the quality of insurance decisions. We must focus behavioral attention far more on the quality of insurance decisions than on the quantity of insurance purchases.

ACKNOWLEDGMENTS

This study project was funded by the Center for the Economic Analysis of Risk.

ETHICS STATEMENT

Ethical approval was granted by the Institutional Review Board of Georgia State University. We are grateful for constructive comments from referees and seminar participants.

REFERENCES

- Belissa, T., Bulte, E., Cecchi, F., Gangopadhyay, S., & Lensink, R. (2019). Liquidity constraints, informal institutions, and the adoption of weather insurance: A randomized controlled trial in Ethiopia. *Journal of Development Economics*, 140, 269–278.
- Carter, M. R., Barrett, C. B., Boucher, S., Chantarat, S., Galarza, F., McPeak, J., Mude, A. G., & Trivelli, C. (2008). Insuring the never before insured: Explaining index insurance through financial education games. *BASIS Briefs BB2008-07*, Department of Agricultural and Applied Economics, University of Wisconsin. <http://www.basis/wisc.edu>
- Carter, M. R., & Chiu, T. (2018). Quality standards for agricultural index insurance: An agenda for action. *The State of Microinsurance, Issue #4, 2018*. The Microinsurance Network. <https://microinsurancenetwork.org/state-microinsurance>
- Casaburi, L., & Willis, J. (2018). Time versus state in insurance: Experimental evidence from contract farming in Kenya. *American Economic Review*, 108, 3778–3813.
- Ceballos, F., & Robles, M. (2020). Demand heterogeneity for index-based insurance: The case for flexible products. *Journal of Development Economics*, 146, 102515.
- Clarke, D. J. (2016). A theory of rational demand for index insurance. *American Economic Journal: Microeconomics*, 8(1), 283–306.
- Cole, S., Giné, X., Tobacman, J., Topalova, P., Townsend, R., & Vickery, J. (2013). Barriers to household risk management: Evidence from India. *American Economic Journal: Applied Economics*, 5(1), 104–135.
- Cole, S., Stein, D., & Tobacman, J. (2014). Dynamics of demand for index insurance: Evidence from a long-run field experiment. *American Economic Review (Papers & Proceedings)*, 104(5), 284–290.
- Eeckhoudt, L., & Schlesinger, H. (2006). Putting risk in its proper place. *American Economic Review*, 96, 280–289.

- Ericson, K. M., & Sydnor, J. (2017). The questionable value of having a choice of levels of health insurance coverage. *Journal of Economic Perspectives*, 31(4), 51–72.
- Fischbacher, U. (2007). z-Tree: zurich toolbox for ready-made economic experiments. *Experimental Economics*, 10(2), 171–178.
- Flatnes, J. E., Carter, M. R., & Mercovich, R. (2018). Improving the quality of index insurance with a satellite-based conditional audit contract, *Unpublished Manuscript*. Department of Agricultural & Resource Economics. Davis: University of California.
- Gao, X. S., Harrison, G. W., & Tchernis, R. (2023). Behavioral welfare economics and risk preferences: A Bayesian approach. *Experimental Economics*, 26, 273–303.
- Gaurav, S., Cole, S., & Tobacman, J. (2011). Marketing complex financial products in emerging markets: Evidence from rainfall insurance in India. *Journal of Marketing Research*, 48(SPL), S150–S162.
- Harrison, G., Morsink, K., & Schneider, M. (2022). Literacy and the quality of index insurance decisions. *The Geneva Risk and Insurance Review*, 47(1), 66–97.
- Harrison, G. W. (2011). Experimental methods and the welfare evaluation of policy lotteries. *European Review of Agricultural Economics*, 38(3), 335–360.
- Harrison, G. W. (2019). The behavioral welfare economics of insurance. *Geneva Risk & Insurance Review*, 44(2), 137–175.
- Harrison, G. W. (2024). Risk preferences and risk perceptions in insurance experiments: Some methodological challenges. *Geneva Risk & Insurance Review*, 49, 127–161.
- Harrison, G. W. (2025). The end of behavioral insurance. In G. Dionne (Ed.), *Handbook of insurance* (3rd ed.). Springer.
- Harrison, G. W., Martínez-Correa, J., & Swarthout, J. T. (2015). Reduction of compound lotteries with objective probabilities: Theory and evidence. *Journal of Economic Behavior & Organization*, 119, 32–55.
- Harrison, G. W., & Ng, J. M. (2016). Evaluating the expected welfare gain from insurance. *Journal of Risk and Insurance*, 83(1), 91–120.
- Harrison, G. W., & Swarthout, J. T. (2023). Cumulative prospect theory in the laboratory: A reconsideration. In G. W. Harrison, & D. Ross (Eds.), *Models of risk preferences: descriptive and normative challenges*. Emerald, Research in Experimental Economics.
- Hey, J. D., & Orme, C. (1994). Investigating generalizations of expected utility theory using experimental data. *Econometrica*, 62(6), 1291–1326.
- Imbens, G. W., & Rubin, D. B. (2015). *Causal inference in statistics, social, and biomedical sciences*. Cambridge University Press.
- Jaspersen, J. G., Ragin, M. A., & Sydnor, J. R. (2022). Predicting insurance demand from risk attitudes. *Journal of Risk and Insurance*, 89, 63–96.
- Loomes, G., & Sugden, R. (1998). Testing different stochastic specifications of risky choice. *Economica*, 65(260), 581–598.
- Mobarak, A. M., & Rosenzweig, M. R. (2013). Informal risk sharing, index insurance, and risk taking in developing countries. *American Economic Review (Papers & Proceedings)*, 103(3), 375–380.
- Monroe, B. (2023). The welfare consequences of individual-level risk preference estimation. In G. W. Harrison, & D. Ross (eds.), *Models of risk preferences: Descriptive and normative challenges*. Emerald, Research in Experimental Economics.
- Norton, M., Osgood, D., Madajewicz, M., Holthaus, E., Peterson, N., Diro, R., Mullally, C., Teh, T.-L., & Gebremichael, M. (2014). Evidence of demand for index insurance: Experimental games and commercial transactions in Ethiopia. *Journal of Development Studies*, 50(5), 630–648.
- Quiggin, J. (1993). *Generalized expected utility theory: The rank-dependent model*. Kluwer Academic.
- Segal, U. (1987). The ellsberg paradox and risk aversion: An anticipated utility approach. *International Economic Review*, 28, 175.
- Segal, U. (1988). Probabilistic insurance and anticipated utility. *The Journal of Risk and Insurance*, 55, 287–297.
- Segal, U. (1990). Two-stage lotteries without the reduction axiom. *Econometrica*, 58(2), 349–377.
- Takahashi, K., Ikegami, M., Sheahan, M., & Barrett, C. B. (2016). Experimental evidence on the drivers of index-based livestock insurance demand in Southern Ethiopia. *World Development*, 78, 324–340.
- Wilcox, N. T. (2023). Unusual estimates of probability weighting functions. In Harrison, G. W. and Ross, D., (Eds.), *Models of risk preferences: Descriptive and normative challenges*, 69–106. Emerald.

SUPPORTING INFORMATION

Additional supporting information can be found online in the Supporting Information section at the end of this article.

How to cite this article: Harrison, G., Martínez-Correa, J., Morsink, K., Ng, J. M., & Swarthout, J. T. (2025). Welfare consequences of the compound risks of index insurance. *Journal of Risk and Insurance*, 1–31. <https://doi.org/10.1111/jori.70012>