



Self-protection, insurance demand and cost-sharing strategy under mean-variance preferences

Wentao Hu, Jan Dhaene & Yiyang Zhang

To cite this article: Wentao Hu, Jan Dhaene & Yiyang Zhang (28 Jan 2026): Self-protection, insurance demand and cost-sharing strategy under mean-variance preferences, Scandinavian Actuarial Journal, DOI: [10.1080/03461238.2026.2622432](https://doi.org/10.1080/03461238.2026.2622432)

To link to this article: <https://doi.org/10.1080/03461238.2026.2622432>



Published online: 28 Jan 2026.



Submit your article to this journal [↗](#)



Article views: 77



View related articles [↗](#)



View Crossmark data [↗](#)



Self-protection, insurance demand and cost-sharing strategy under mean-variance preferences

Wentao Hu^a, Jan Dhaene^b and Yiyang Zhang^c

^aSchool of Finance, Renmin University of China, Beijing, People's Republic of China; ^bFaculty of Economics and Business, KU Leuven, Leuven, Belgium; ^cThe corresponding author. Department of Mathematics, Southern University of Science and Technology, Shenzhen, Guangdong Province, People's Republic of China

ABSTRACT

This article studies a mean-variance Stackelberg game between an insured and an insurer under a self-protection model, where the insured's effort is observable by the insurer and reflected in the insurance premium. We adopt a proportional insurance model under a quadratic premium principle and incorporate a cost-sharing mechanism for self-protection. The existence of equilibrium solutions for insurance demand, self-protection effort, and cost-sharing proportions is established, and their interactions are analyzed theoretically. Numerical analysis reveals conditions under which self-protection and market insurance act as complementary or substitutable risk management tools.

ARTICLE HISTORY

Received 15 March 2025
Accepted 22 January 2026

KEYWORDS

Self-protection; proportional insurance; mean-variance preferences; quadratic premium principles; cost-sharing

JEL CLASSIFICATIONS

C61; G22; G32

1. Introduction

In risk management, a central objective is to encourage insureds to adjust their behavior in ways that reduce exposure to risk. Self-protection, a key mechanism in this regard, refers to preventive efforts aimed at lowering the probability of a loss. The seminal study by Ehrlich and Becker (1972) established the theoretical foundation for analyzing such behavior, showing that self-protection and market insurance can be complementary as the premium loading changes when effort is fully reflected in the premium. When deciding whether and to what extent to engage in self-protection, individuals weigh its costs against its benefits. Knowing how these decisions are made is important for designing policies and incentives that promote self-protection, benefiting both insurers and insureds. For related studies on self-protection decision-making, see Chiu (2005), Meyer and Meyer (2011), Peter (2017), L. Li and Peter (2021), Crainich and Menegatti (2021), Augeraud-Véron and Leandri (2024), Yin and Meng (2025) and J. Li et al. (2025). Besides, the influence of risk preferences on prevention decisions is well-documented. For instance, Peter (2021a, 2021c) have examined its role in shaping optimal strategies. Further research explores specific facets such as loss aversion (Macé & Peter, 2021) and prudence (Chi et al., 2025) in intertemporal models. Interested readers can refer to Peter (2024b) for more details, which provides a selective survey of the economics of self-protection.

When self-protection is considered within the framework of insurance pricing, taking preventive efforts can potentially reduce insurance premiums, creating a financial incentive for prevention. This phenomenon has been studied extensively. Using the Dual Theory of Choice, Courbage (2001) argued that without proper pricing mechanisms, insureds may favor self-protection over insurance,

weakening the risk-sharing function of insurance markets. Building on this, Kelly and Kleffner (2003) examined a case where a risk-neutral monopolistic insurer sets premiums based on the price elasticity of demand for coverage. Extending these insights, Bensalem et al. (2020) developed an equilibrium model in which the insurer sets the premium and the potential insured determines how much insurance to purchase and how much to invest in self-protection to reduce the risk. Following this framework, Zeller and Scherer (2023) proposed a Stackelberg game in which both insurer and insured use distortion risk measures. In the context of peer-to-peer insurance, Z. Chen, Dhaene, et al. (2024) studied optimal self-protection while accounting for social connections. For further studies on optimal decision-making in self-protection and insurance demand, see Courbage and Coulon (2004), Zweifel et al. (2012), Menegatti (2014), Hofmann and Rothschild (2019) and L. Chen et al. (2026).

In practice, insurers often use cost-sharing mechanisms to incentivize insureds to maintain or increase preventive efforts, rather than rely solely on insurance compensation, thereby reducing the insurer's risk exposure. For example, in health insurance, insurers may provide free smart devices or dynamic health management programs to enable timely interventions, lowering both the likelihood and severity of diseases (Khan & Yairi, 2018). For epidemic risks, insurers can offer free precautionary packages to high-risk policyholders (Cui et al., 2021). Similarly, as noted in Zeller and Scherer (2023), insurers may provide more affordable cybersecurity services to enhance protection against cyber threats. In general, the shared costs of self-protection constitute an expense for the insurer. Consequently, insurers must carefully balance the costs of supporting self-protection against the potential reductions in their own risk exposure.

In this article, we model a mean-variance Stackelberg game between an insured and an insurer under the self-protection model. The insured exerts preventive effort, which is fully reflected in the premium pricing and thus reduces the insurance premium. Within this framework, we consider proportional insurance under a quadratic premium principle,¹ combined with a cost-sharing mechanism for self-protection. First, the insured chooses the optimal insurance coverage and preventive effort to maximize a mean-variance functional of their wealth. Then, the insurer selects the optimal cost-sharing proportion to maximize its expected wealth.

The main contributions of this article are as follows:

First, our work provides theoretical results for a mean-variance Stackelberg game with cost-sharing and a quadratic convex premium principle. We establish the existence of optimal solutions for insurance demand, self-protection effort, and the cost-sharing ratio, and analyze their variations under different conditions. Explicit expressions for the optimal insurance proportion and protection level are derived. In particular, the optimal insurance demand can take partial insurance, rather than being confined to corner solutions such as null or full insurance.

Second, through numerical analysis, we examine how the relationship between market insurance and self-protection switches between substitution and complementarity as critical factors change, thereby contributing to the literature on the interaction between insurance demand and preventive efforts. Our results help design insurance mechanisms that promote insureds' prevention efforts while maximizing the interests of both parties. Such mechanisms are key to encouraging proactive risk management and improving market resilience.

The remainder of the paper is organized as follows. Section 2 introduces the basic setting and the optimization model. Section 3 analyzes the insured's problem and derives the optimal insurance demand and self-protection effort. Section 4 studies the insurer's problem and determines the optimal cost-sharing coefficient. Section 5 presents numerical results. Section 6 concludes.

¹ Unlike the expected-value premium principle, the quadratic premium principle, one example of convex premium principles, captures the benefits of risk pooling, showing that diversifying risks across multiple independent policies can lower the total premium. For details, see Section 2 and Deprez and Gerber (1985), Kaluszka (2005) and Cao et al. (2024a).

2. Model setting

In this section, we present the main model setting. We begin by outlining the basic setup, followed by a discussion of self-protection, the costs of preventive efforts, and the cost-sharing arrangement. Finally, we formulate the mean-variance Stackelberg game that captures the interaction between the insured and the insurer.

2.1. Basic setting

Throughout this work, all random variables are defined on a common probability space $(\Omega, \mathcal{F}, \mathbb{P})$. Let X be a non-negative random variable on $\mathbb{R}_+ = [0, +\infty)$, representing the aggregate loss faced by the insured. In this paper, we focus on proportional insurance, in which the indemnity function takes the form $I(x) = ax$, with $0 \leq a \leq 1$. Compared to full insurance, this contract requires the insured to bear a portion of the total loss when a claim arises, which contributes to the policyholder's risk awareness (He et al., 2008; Wu & Colwell, 1988; Zeller & Scherer, 2023).

In exchange for covering the ceded loss $I(X)$, the insurer charges a premium determined by the quadratic premium principle:

$$\pi(I(X)) = \mathbb{E}[h(I(X))], \quad (1)$$

where

$$h(x) = \frac{\theta_2}{2}x^2 + \theta_1x, \quad \theta_2 \geq 0, \quad \theta_1 \geq 1. \quad (2)$$

The quadratic premium principle belongs to the class of convex premium principles.² It has become a common tool in dynamic models of optimal insurance (Cao et al., 2023, 2024b). The expected-value premium principle arises as a special case when $\theta_2 = 0$. Although the quadratic premium and mean-variance premium principles may look similar under certain parameter choices, they are not equivalent. In particular, $\mathbb{E}[h(\cdot)]$ exhibits convexity, whereas the mean-variance premium lacks this property (Klimaviciute, 2017).

2.2. Self-protection

We now introduce the framework of self-protection, wherein the insured can exert preventive effort to reduce the probability of a potential loss. Denote the effort by $e \in \mathbb{R}_+$. The resulting loss is described by the non-negative random variable X_e , whose distribution is given by

$$\mathbb{P}_{X_e} = (1 - p(e))\delta_{\{0\}} + p(e)\mathbb{P}_Y,$$

where $0 < p(e) < 1$, $\delta_{\{0\}}$ is the Dirac mass at zero, and $\mathbb{P}_Y$ is the distribution of a positive continuous random variable. Thus, X_e equals Y with probability $p(e)$, and equals zero with probability $1 - p(e)$. We assume $p'(e) < 0$ and $p''(e) > 0$ (Courbage, 2001; Courbage et al., 2017; Crainich & Eeckhoudt, 2017; Hofmann & Peter, 2015), meaning that greater effort reduces the probability of loss, while the marginal effect of additional effort is decreasing. We further assume $0 < \underline{p} = p(+\infty)$ and $p(0) = \bar{p} < 1$.

We assume that prevention effort is fully observable by the insurer and reflected in premium calculation. After exerting effort e , the insured faces loss X_e and receives coverage $I(X_e) = aX_e$ in exchange for a premium:

$$\mathbb{E}[h(I(X_e))] = \frac{\theta_2}{2}a^2\mathbb{E}^2[X_e] + \theta_1a\mathbb{E}[X_e]. \quad (3)$$

² The convex premium principle is first introduced by Deprez and Gerber (1985) as an extension of the classical premium principle. Kaluszka (2005) and Cao et al. (2024a) studied optimal reinsurance and insurance design within this framework.

2.3. Cost of effort

In practice, prevention effort can often be interpreted as services that are directly available in the market. For example, for cyber risks, firms may purchase cyber-security services which can lower the probability of data breaches and improve resilience against cyberattacks (Chase et al., 2017; Lin et al., 2009; Ram & Sreenivaasan, 2010). For epidemic risks, vaccination programs reduce both the probability of infection and the expected severity of illness (Z. Chen, Chen, et al., 2024; Hu et al., 2024; Williams et al., 2021). For health risks, individuals may invest in preventive healthcare (Huang et al., 2017; Madubuike et al., 2024; Wang et al., 2016).

Inspired by these examples, we denote $c(e)$ as the cost of prevention effort e , which is increasing and convex, satisfying $c(0) = 0$, $c'(e) > 0$, and $c''(e) > 0$. The condition $c''(e) > 0$ captures the feature of increasing marginal costs of self-protection: while initial safety measures are relatively easy to implement, further risk reduction typically requires more sophisticated and costly technologies. The model configuration is based on the existing literature; see, for example, Courbage (2001), Meyer and Meyer (2011), Denuit et al. (2016), Peter (2017) and Zeller and Scherer (2023).

In this paper, we assume that the insured can access self-protection services through the insurer, who enjoys a lower price when purchasing these services. This assumption is reasonable because an insurer, by aggregating the demand of many insureds, can purchase in bulk. Such bulk purchasing power allows the insurer to secure quantity discounts from suppliers, who are willing to offer better terms in exchange for higher volume orders. In contrast, an individual insured's purchase is typically at the retail level and lacks such negotiating leverage. This setting reflects economies of scale, an effect well documented in the literature (see, e.g. Altintas et al., 2008; Corbett & De Groote, 2000; Lee & Rosenblatt, 1986; Monahan, 1984; Weng, 1995; Zeller & Scherer, 2023).

Accordingly, we assume that the cost of effort e for the insurer is given by $c(e)$, while that for the insured is $\beta_0 c(e)$ with $\beta_0 > 1$.

2.4. Cost-sharing mechanism

We now introduce the cost-sharing arrangement. Under this arrangement, the insured's cost consists of two parts. First, based on the insurance coverage ratio a , the insured can access the corresponding share of self-protection services at the wholesale price $c(e)$. For this part, the insurer shares a proportion of the effort cost, and the insured bears the remaining fraction $\beta \in [0, 1]$. The insured's cost in this component, $C_{b_1}(e)$, is therefore

$$C_{b_1}(e) = a\beta c(e). \quad (4)$$

Second, for the uninsured share $(1 - a)$, the insured must purchase effort directly from suppliers at the retail price $\beta_0 c(e)$. The corresponding cost, $C_{b_2}(e)$, is given by

$$C_{b_2}(e) = (1 - a)\beta_0 c(e). \quad (5)$$

Summing these two parts, the total cost borne by the insured, $C_b(e)$, can be expressed as

$$C_b(e) = [\beta_0 - a(\beta_0 - \beta)]c(e). \quad (6)$$

From (4), it follows that the insurer's cost share, $C_s(e)$, is

$$C_s(e) = a(1 - \beta)c(e). \quad (7)$$

When the insured do not purchase insurance ($a = 0$), they cannot access any preventive effort through the insurer and must bear the full cost at the retail price. In this case, $C_b(e) = \beta_0 c(e)$. Accordingly, $C_s(e) = 0$. When the insured purchase full insurance ($a = 1$), they can obtain all preventive effort at the wholesale price through the insurer. Hence, $C_b(e) = \beta c(e)$. The insurer covers the remaining cost, $C_s(e) = (1 - \beta)c(e)$.

If the insurer does not share any cost ($\beta = 1$), then $C_s(e) = 0$ and the insured only benefit from the wholesale price discount. In this case, $C_b(e) = [\beta_0 - a(\beta_0 - 1)]c(e)$. Conversely, if the insurer fully covers the cost associated with the coverage proportion ($\beta = 0$), then $C_s(e) = ac(e)$, and the insured pays $C_b(e) = \beta_0(1 - a)c(e)$ for the remaining effort.

This cost-sharing mechanism has a practical foundation. First, the mechanism adopted in this paper follows Zeller and Scherer (2023), which is based on real-world observations of cyber risk. A similar structure is used in Cui et al. (2021), where, in the context of epidemic risk, insurers allocate part of the premium income to subsidize policyholders' preventive actions, such as vaccination, in order to maximize their own benefits. In addition, the Shaanxi province has a policy in place that insurers must invest in accident prevention services, allocating between 18% and 21% of insurance premiums (Shaanxi, 2025). More discussion and results on the cost-sharing mechanism can be found in Appendix A.2.

2.5. Mean-variance Stackelberg game

In this section, we describe a Stackelberg game between the insurer and the insured. The insured acts as the follower, whose decision variables are the insurance demand a and the prevention effort e . The insurer is the leader with commitment power and can act as a monopolist, thereby preventing competitors from offering marginally better cost-sharing arrangements³. This allows the insurer to set a profit-maximizing cost-sharing rate β without facing competition.

The objective of the insured is to maximize a mean-variance functional of terminal wealth given β . Let the initial wealth of the insured be w_b . The insured's terminal wealth W_b is given by

$$\begin{aligned} W_b &= w_b - \mathbb{E}[h(I(X_e))] - X_e + I(X_e) - C_b(e) \\ &= w_b - \frac{\theta_2}{2} a^2 \mathbb{E}^2[X_e] - \theta_1 a \mathbb{E}[X_e] - (1 - a)X_e - C_b(e). \end{aligned}$$

The insured then solves the following optimization problem⁴:

Problem 2.1

$$\sup_{(e,a) \in \mathbb{R}_+ \times [0,1]} \left\{ \mathbb{E}[W_b] - \frac{\gamma_1}{2} \text{Var}[W_b] \right\}, \quad (8)$$

where γ_1 measures the insured's risk aversion.

The objective of the insurer is to maximize its expected terminal wealth, taking into account the optimal responses $a^*(\beta)$ and $e^*(\beta)$ from the insured. Let the initial wealth of the insurer be w_s . The insurer's terminal wealth W_s is

$$\begin{aligned} W_s &= w_s + \mathbb{E}[h(I(X_{e^*(\beta)}))] - I(X_{e^*(\beta)}) - C_s(e^*(\beta)) \\ &= w_s + \frac{\theta_2}{2} a^*(\beta)^2 \mathbb{E}^2[X_{e^*(\beta)}] + \theta_1 a^*(\beta) \mathbb{E}[X_{e^*(\beta)}] - a^*(\beta) X_{e^*(\beta)} - C_s(e^*(\beta)). \end{aligned}$$

Accordingly, the insurer solves:

Problem 2.2

$$\max_{\beta \in [0,1]} \mathbb{E}[W_s]. \quad (9)$$

³ Insurers often act as both a leader and a monopolist. First, in the actuarial literature on Stackelberg games, the insurer is commonly modeled as the leader. Second, in state-monopolized or specialized markets, such as China's compulsory traffic insurance or nuclear insurance, one or very few insurers often dominate. Furthermore, in environments where traditional sales channels dominate and information is not transparent, policyholders struggle to compare prices and terms. This reinforces the insurer's position as the absolute leader in product design and pricing, effectively granting it an oligopolistic role.

⁴ Under the mean-variance objective, the optimal insurance contract is a changing stop-loss. However, due to its mathematical complexity, we instead consider a proportional insurance contract. Details are given in Appendix A.1.

In the following two sections, we solve the model by backward induction, starting with the insured's problem, followed by the insurer's.

3. Solution to the insured's problem

In this section, we solve the insured's Problem 2.1, which is equivalent to

$$\inf_{(e,a) \in \mathbb{R}_+ \times [0,1]} L(e, a), \quad (10)$$

where

$$L(e, a) = (1 - a)\mathbb{E}[X_e] + \frac{\theta_2}{2}a^2\mathbb{E}[X_e]^2 + \theta_1a\mathbb{E}[X_e] + [\beta_0 - a(\beta_0 - \beta)]c(e) + \frac{\gamma_1}{2}\text{Var}[(1 - a)X_e].$$

We first fix e and solve for the optimal $a^*(e)$ that minimizes $L(e, a)$, and then determine e^* based on $a^*(e)$.

3.1. Optimal insurance demand under given effort

For any given e , the sub-optimization problem can be given by:

$$\inf_{a \in [0,1]} L_e(a), \quad (11)$$

where

$$L_e(a) = (1 - a)\mathbb{E}[X_e] + \frac{\theta_2 a^2}{2} \mathbb{E}[X_e]^2 + \theta_1 a \mathbb{E}[X_e] - a(\beta_0 - \beta)c(e) + \frac{\gamma_1}{2} \text{Var}[(1 - a)X_e].$$

In the sake of convenience, we denote $\mu_0 = \mathbb{E}[Y]$, $\sigma_0^2 = \text{Var}[Y]$, and $\lambda_0 = \mathbb{E}[Y^2]$. Define two functions of e as follows:

$$\begin{aligned} f_1(e) &= \gamma_1 \left(\frac{\lambda_0}{\mu_0} - p(e)\mu_0 \right) + (\beta_0 - \beta) \frac{c(e)}{p(e)\mu_0}, \\ f_2(e) &= -\theta_2 p(e)\mu_0 + (\beta_0 - \beta) \frac{c(e)}{p(e)\mu_0}. \end{aligned}$$

To establish the main results, we first need the following lemma:

Lemma 3.1: *The following statements hold:*

- Case (i): If $\theta_1 - 1 \geq \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right)$, then there exists a unique $0 \leq \underline{e} < \bar{e}$ such that $\theta_1 - 1 = f_1(\underline{e})$ and $\theta_1 - 1 = f_2(\bar{e})$.
- Case (ii): If $\theta_1 - 1 < \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right)$, then there exists a unique $0 < \tilde{e}$ such that $\theta_1 - 1 = f_2(\tilde{e})$.

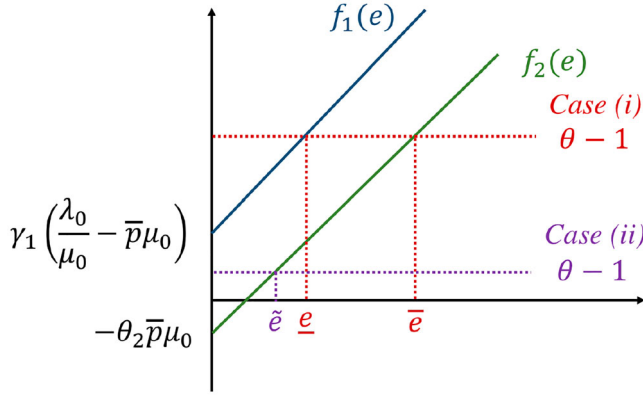


Figure 1. Illustration of Case (i) and Case (ii).

Proof: The approach for determining $\underline{e}$, $\bar{e}$, and $\tilde{e}$ is illustrated in Figure 1. Firstly, given the inequality

$$\frac{\lambda_0}{\mu_0} - p(e)\mu_0 = \frac{\text{Var}[X_e]}{\mathbb{E}[X_e]} > 0,$$

we can conclude that for all $e \in \mathbb{R}^+$, $f_1(e) \geq f_2(e)$. Secondly, since $p(e) > 0$, $p'(e) < 0$, $c(e) > 0$, and $c'(e) > 0$, we also know that $f_1'(e) > 0$ and $f_2'(e) > 0$. This implies the following results:

$$\begin{aligned} \min f_1(e) &= f_1(0) = \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right), & \min f_2(e) &= f_2(0) = -\theta_2 \bar{p}\mu_0, \\ \max f_1(e) &= f_1(+\infty) = +\infty, & \max f_2(e) &= f_2(+\infty) = +\infty. \end{aligned}$$

Thus, if $\theta_1 - 1 \geq \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right)$, there exists a unique pair $0 \leq \underline{e} < \bar{e}$ such that

$$\theta_1 - 1 = f_1(\underline{e}) \quad \text{and} \quad \theta_1 - 1 = f_2(\bar{e}).$$

Alternatively, if $\theta_1 - 1 < \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right)$, then there uniquely exists $0 < \tilde{e}$ such that

$$\theta_1 - 1 = f_2(\tilde{e}).$$

Thus, the lemma is proved. ■

By Lemma 3.1, *Case (i)* corresponds to the scenario of a high marginal price of risk, while *Case (ii)* corresponds to a low marginal price. For the sake of clarity, we will specifically refer to these two scenarios as *Case (i)* and *Case (ii)* in the following text. The domain of effort e can be divided into either $[0, \underline{e}]$, $(\underline{e}, \bar{e})$, and $[\bar{e}, +\infty)$ under *Case (i)*, or $[0, \tilde{e})$ and $[\tilde{e}, +\infty)$ under *Case (ii)*, depending on the value of θ_1 , with thresholds $\underline{e}$, $\bar{e}$, $\tilde{e}$ (as illustrated in Figure 1). This partitioning serves as the basis of subsequent analysis. The optimization problem is solved within each interval, and the global solution is obtained by combining these cases.

Based on Lemma 3.1, we now provide the form of the optimal $a^*(e)$ for any given e .

Proposition 3.2: For any given effort $e \in \mathbb{R}_+$, the following conditions hold:

- In *Case (i)*, we obtain:
 - (1) If $e \in [0, \underline{e}]$, then $a^*(e) = 0$;

(2) If $e \in (\underline{e}, \bar{e})$, then $a^*(e) \in (0, 1)$ and $a^*(e)$ is given by:

$$a^*(e) = \frac{(\beta_0 - \beta)c(e) - (\theta_1 - 1)p(e)\mu_0 + \gamma_1 p(e)\lambda_0 - \gamma_1 p^2(e)\mu_0^2}{\theta_2 p^2(e)\mu_0^2 + \gamma_1 p(e)\lambda_0 - \gamma_1 p^2(e)\mu_0^2}; \quad (12)$$

(3) If $e \in [\bar{e}, +\infty)$, then $a^*(e) = 1$.

• In Case (ii), we obtain:

(1) If $e \in [0, \tilde{e})$, then $a^*(e) \in (0, 1)$ and is given by Equation (12);

(2) If $e \in [\tilde{e}, +\infty)$, then $a^*(e) = 1$.

Proof: We begin by calculating the first and second derivatives of $L_e(a)$ with respect to a . The results are as follows:

$$\begin{aligned} L'_e(a) &= (\theta_1 - 1)p(e)\mu_0 + \theta_2 a p^2(e)\mu_0^2 - \gamma_1(1 - a)(p(e)\lambda_0 - p^2(e)\mu_0^2) - (\beta_0 - \beta)c(e), \\ L''_e(a) &= \theta_2 p^2(e)\mu_0^2 + \gamma_1(p(e)\lambda_0 - p^2(e)\mu_0^2), \end{aligned}$$

which implies $L''_e(a) > 0$, meaning $L_e(a)$ is convex in a . Next, we calculate the minimum and maximum values of $L'_e(a)$:

$$\begin{aligned} \min L'_e(a) &= L'_e(0) = (\theta_1 - 1)p(e)\mu_0 - \gamma_1(p(e)\lambda_0 - p^2(e)\mu_0^2) - (\beta_0 - \beta)c(e), \\ \max L'_e(a) &= L'_e(1) = (\theta_1 - 1)p(e)\mu_0 + \theta_2 p^2(e)\mu_0^2 - (\beta_0 - \beta)c(e). \end{aligned}$$

Based on this, we can determine the optimal solution $a^*(e)$ as follows:

- (1) If $\min L'_e(a) \geq 0$, which is equivalent to $\theta_1 - 1 \geq f_1(e)$, then $L'_e(a) \geq 0$, and the optimal value is $a^*(e) = 0$.
- (2) If $\max L'_e(a) \leq 0$, which is equivalent to $\theta_1 - 1 \leq f_2(e)$, then $L'_e(a) \leq 0$, and the optimal value is $a^*(e) = 1$.
- (3) If $\min L'_e(a) < 0$ and $\max L'_e(a) > 0$, which is equivalent to $f_2(e) < \theta_1 - 1 < f_1(e)$, then there exists a unique solution $a^*(e) \in (0, 1)$, given by:

$$\begin{aligned} a^*(e) &= \frac{(\beta_0 - \beta)c(e) - (\theta_1 - 1)\mathbb{E}[X_e] + \gamma_1 \text{Var}(X_e)}{\theta_2 \mathbb{E}[X_e]^2 + \gamma_1 \text{Var}(X_e)} \\ &= \frac{(\beta_0 - \beta)c(e) - (\theta_1 - 1)p(e)\mu_0 + \gamma_1 p(e)\lambda_0 - \gamma_1 p^2(e)\mu_0^2}{\theta_2 p^2(e)\mu_0^2 + \gamma_1 p(e)\lambda_0 - \gamma_1 p^2(e)\mu_0^2}, \end{aligned}$$

such that $L'_e(a) \leq 0$ when $a \in [0, a^*(e)]$ and $L'_e(a) \geq 0$ when $a \in [a^*(e), 1]$.

Based on Lemma 3.1, we know that $f_2(e) \leq f_1(e)$, with:

$$\begin{aligned} \min f_1(e) &= \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right), \quad \max f_1(e) = +\infty, \\ \min f_2(e) &= -\theta_2 \bar{p}\mu_0, \quad \max f_2(e) = +\infty. \end{aligned}$$

Thus, the following cases arise:

- When $\theta_1 - 1 \geq \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right)$:
 - (1) For $e \in [0, \underline{e}]$, $\theta_1 - 1 \geq f_1(e)$ and $a^*(e) = 0$.
 - (2) For $e \in (\underline{e}, \bar{e})$, $f_2(e) < \theta_1 - 1 < f_1(e)$ and $a^*(e)$ is given by (12).
 - (3) For $e \in [\bar{e}, +\infty)$, $\theta_1 - 1 \leq f_2(e)$ and $a^*(e) = 1$.

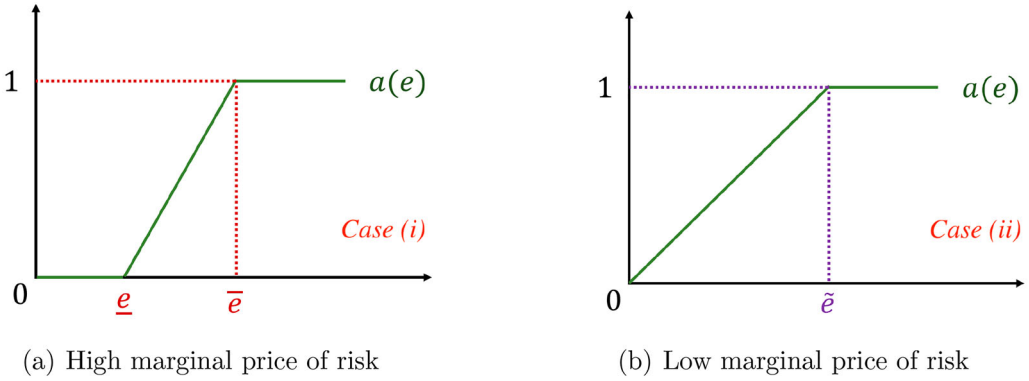


Figure 2. Plot of $a^*(e)$ with respect to e . (a) High marginal price of risk and (b) Low marginal price of risk.

- When $\theta_1 - 1 < \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right)$:
 - (1) For $e \in [0, \tilde{e})$, $f_2(e) < \theta_1 - 1 < f_1(e)$ and $a^*(e)$ is given by (12).
 - (2) For $e \in [\tilde{e}, +\infty)$, $\theta_1 - 1 \leq f_2(e)$ and $a^*(e) = 1$.

In conclusion, the infimum of the objective function is achieved at its minimum value, which implies that $a^*(e)$ is the solution to:

$$\inf_{a \in [0,1]} L_e(a) = \min_{a \in [0,1]} L_e(a).$$

Hence, the proof is finished. ■

Proposition 3.2 shows that when the marginal price of risk is relatively high, as illustrated in Figure 2(a), the optimal insurance choice exhibits three regimes:

- (1) For low effort levels (i.e. $e \leq \underline{e}$), no insurance is purchased (i.e. $a^*(e) = 0$);
- (2) For intermediate effort (i.e. $\underline{e} < e < \bar{e}$), proportional insurance is chosen (i.e. $0 < a^*(e) < 1$);
- (3) For high effort (i.e. $e \geq \bar{e}$), full insurance is obtained (i.e. $a^*(e) = 1$).

This pattern indicates a complementary relationship between self-protection and insurance demand, since greater preventive effort is associated with higher coverage.⁵ The complementarity discussed here concerns the joint behavior of e and $a^*(e)$, rather than the equilibrium pair $(e^*, a^*(e^*))$.

When the marginal price of risk is relatively low, as shown in Figure 2(b), a similar pattern arises: if $e \geq \tilde{e}$ the insured opts for full insurance, while if $e < \tilde{e}$ proportional insurance is chosen. The difference from the previous case is that here $a^*(e)$ never collapses to zero, since the lower marginal price of risk ensures a persistent incentive to buy insurance.

Finally, the thresholds $\tilde{e}$ and $\bar{e}$ in these two cases are both determined by the condition $\theta_1 - 1 = f_2(e)$, which implies that the interior solution for $a^*(e)$ given by Equation (12) in *Case (i)* when $e \in (\underline{e}, \bar{e})$ corresponds exactly to the interior solution in *Case (ii)* when $e \in [0, \tilde{e})$. Likewise, the boundary solution $a^*(e) = 1$ in *Case (i)* for $e \in [\bar{e}, +\infty)$ corresponds to the boundary solution in *Case (ii)* for $e \in [\tilde{e}, +\infty)$.

⁵ Although this complementarity cannot be established analytically, Figure 3 in Section 5 provides supporting numerical analysis evidence.

For any given e , let us define

$$\hat{a}^*(e) = \frac{(\beta_0 - \beta)c(e) - (\theta_1 - 1)p(e)\mu_0 + \gamma_1 p(e)\lambda_0 - \gamma_1 p^2(e)\mu_0^2}{\theta_2 p^2(e)\mu_0^2 + \gamma_1 p(e)\lambda_0 - \gamma_1 p^2(e)\mu_0^2}.$$

Then, the following expression presents a compact way to denote the optimal insurance demand $a^*(e)$ identified in Proposition 3.2:

- In *Case (i)*, we can obtain:

$$a^*(e) = \mathbf{1}_{[\bar{e}, \infty)}(e) + \hat{a}^*(e)\mathbf{1}_{(e, \bar{e})}(e).$$

- In *Case (ii)*, we can obtain:

$$a^*(e) = \mathbf{1}_{[\bar{e}, \infty)}(e) + \hat{a}^*(e)\mathbf{1}_{[0, \bar{e})}(e).$$

The findings in Proposition 3.2 are related to Bensalem et al. (2020), Zeller and Scherer (2023), Q. Li et al. (2025), and S. Chen et al. (2025), but also exhibit important differences. Bensalem et al. (2020) considered a linear premium principle and employed a distortion risk measure as the objective function. Since distortion risk measures are positively homogeneous, the resulting objective function is linear in a , leading to corner solutions where a^* is either 0 or 1. In contrast, our framework incorporates a convex premium function with a variance term, which makes the objective function concave in a and allows interior solutions with $0 < a^* < 1$. Moreover, cost-sharing mechanisms were not considered in Bensalem et al. (2020). Extending that work, Zeller and Scherer (2023) introduced a cost-sharing arrangement but still relied on distortion risk measures and a homogeneous premium principle, and therefore also concluded that a^* must be 0 or 1.

Q. Li et al. (2025) further generalized the distortion-risk-measure framework to convex premium principles, under which partial insurance may be optimal in some cases. Our study differs from Q. Li et al. (2025) in several respects: we adopt mean-variance preferences instead of distortion risk measures, and we explicitly incorporate the insurer's profit through an optimal cost-sharing strategy. We consider a Stackelberg game between the insured and the insurer, while the work of Q. Li et al. (2025) only studied the optimization problem from the perspective of the insured. Finally, S. Chen et al. (2025) examined optimal effort and indemnity design within distortion risk measures without restricting the functional form of insurance policies. More discussion and results can be found in Appendix A.2.

3.2. Optimal prevention effort

In this section, we will discuss e^* based on the expression for $a^*(e)$ derived in the previous section. The optimization problem can be represented as:

$$\inf_{e \in [0, +\infty)} L_{a^*}(e), \quad (13)$$

where

$$\begin{aligned} L_{a^*}(e) = & (1 - a^*(e))\mathbb{E}[X_e] + \frac{\theta_2 a^*(e)^2}{2} \mathbb{E}[X_e]^2 + \theta_1 a^*(e)\mathbb{E}[X_e] \\ & + [\beta_0 - a^*(e)(\beta_0 - \beta)]c(e) + \frac{\gamma_1}{2} \text{Var}[(1 - a^*(e))X_e]. \end{aligned}$$

We now introduce an assumption to establish the main result. We define

$$H(e) = \mathbb{E}[X_e] + \frac{\gamma_1}{2} \text{Var}[X_e]. \quad (14)$$

Comparing (14) with the insured's objective function in (10),

$$L(e, a) = (1 - a)\mathbb{E}[X_e] + \frac{\theta_2 a^2}{2}\mathbb{E}[X_e]^2 + \theta_1 a\mathbb{E}[X_e] + [\beta_0 - a(\beta_0 - \beta)]c(e) + \frac{\gamma_1}{2}\text{Var}[(1 - a)X_e],$$

we can see that $H(e)$ corresponds to a special case of $L(e, a)$, namely when market insurance is absent ($a \equiv 0$) and the cost of prevention effort is not included. We now introduce the following assumption:

Assumption 3.1: Assume that $H'(e) < 0$ and $H''(e) > 0$.

Assumption 3.1 states that the objective function in the special case is decreasing and convex function of effort e . The motivation for Assumption 3.1 is to make the result non-trivial and tractable. First, $H'(e) < 0$ guarantees that effort strictly reduces the cost of risk. In the special case without insurance and without effort costs, this ensures that the optimal choice is $e^* > 0$. Without this condition, we would obtain a corner solution $e^* = 0$. In that case, if we consider effort costs and incorporate insurance, the optimal solution will always be equal to 0 unless the incentive effect of insurance is very large. This makes the result trivial and makes it difficult to observe whether the effect of insurance on effort is motivating or disincentivizing. In fact, our assumption shares a similar motivation with the studies in S. Chen et al. (2025) and Q. Li et al. (2025), it differs in that their assumptions are specifically based on distortion risk measures, a consequence of their choice of objective function.

From the perspective for distribution, while effort always reduces expected loss $\mathbb{E}[X_e]$, but its effect on variance $\text{Var}[X_e]$ is ambiguous. An increase in variance may still be favorable if, for example, the worst-case outcome under effort dominates the best-case outcome without effort. Since the mean-variance framework treats all variance as undesirable, the assumption $H'(e) < 0$ enables us to incorporate such cases of 'beneficial' variance. Some examples are provided in Appendix A.3.

In addition, the condition $H''(e) > 0$ reflects the economic intuition of diminishing returns to prevention: while initial increases in effort may substantially reduce losses, the incremental benefit diminishes as effort rises further. Together, these conditions allow us to analyze the interaction between insurance demand and prevention effort in a tractable and economically interpretable way.

We now turn to the solution of problem (13) under Assumption 3.1. The procedure is as follows: Based on the partitions of the domain of e obtained in *Case (i)* and *Case (ii)*, we identify the candidate minimizers $\{e_i^*\}$ in each interval. We then determine the global optimum e^* by selecting among these local minima.

For convenience, let $L_0(e)$, $L_1(e)$, and $L_{a^*}(e)$ denote the objective function $L(e, a)$ evaluated at $a = 0$, $a = 1$, and $a = a^*(e)$, respectively, where $a^*(e)$ is given by (12). Then

- In *Case (i)*, we introduce the following notations:
 - (1) For the case $e \in [0, \underline{e}]$, we define e_0^* such that

$$e_0^* = \arg \inf_{e \in [0, \underline{e}]} L_0(e),$$

- (2) For the case $e \in [\underline{e}, \bar{e}]$, we define e_1^* such that:

$$e_1^* = \arg \inf_{e \in [\underline{e}, \bar{e}]} L_{a^*}(e).$$

- (3) For the case $e \in [\bar{e}, +\infty)$, we define e_2^* such that

$$e_2^* = \arg \inf_{e \in [\bar{e}, +\infty)} L_1(e),$$

- In *Case (ii)*, we introduce the following notations:

(1) For the case $e \in [0, \tilde{e}]$, we define $\tilde{e}_1^*$ such that:

$$\tilde{e}_1^* = \arg \inf_{e \in [0, \tilde{e}]} L_{a^*}(e).$$

(2) For the case $e \in [\tilde{e}, +\infty)$, we define $\tilde{e}_2^*$ such that

$$\tilde{e}_2^* = \arg \inf_{e \in [\tilde{e}, \infty)} L_1(e).$$

Effort level e_0^* only arises in *Case (i)* since only there $a^*(e) = 0$ is possible. In contrast, when $0 < a^*(e) \leq 1$, the minimizers need to be analyzed in both *Case (i)* and *Case (ii)*. We first discuss e_0^* and obtain:

Lemma 3.3: *In Case (i) under Assumption 3.1, we obtain the following:*

- If $\beta_0 \leq -\frac{H'(\underline{e})}{c'(\underline{e})}$, then $e_0^* = \underline{e}$;
- If $-\frac{H'(\underline{e})}{c'(\underline{e})} < \beta_0 < -\frac{H'(0)}{c'(0)}$, then there exists a unique $e_0^* \in (0, \underline{e})$, where e_0^* satisfies

$$H'(e_0^*) + \beta_0 c'(e_0^*) = 0;$$

- If $\beta_0 \geq -\frac{H'(0)}{c'(0)}$, then $e_0^* = 0$.

Proof: We know that

$$L_0(e) = p(e) \left(\mu_0 + \frac{\gamma_1}{2} \lambda_0 \right) - p^2(e) \frac{\gamma_1}{2} \mu_0^2 + \beta_0 c(e).$$

Then, we analyze the derivatives of $L_0(e)$:

$$L'_0(e) = H'(e) + \beta_0 c'(e), \quad L''_0(e) = H''(e) + \beta_0 c''(e) > 0.$$

This shows that $L'_0(e)$ is increasing on the interval $[0, \underline{e}]$. We have:

$$\min L'_0(e) = L'_0(0), \quad \max L'_0(e) = L'_0(\underline{e}).$$

Now, considering the three cases:

- (1) If $L'_0(\underline{e}) \leq 0$, that is $\beta_0 \leq -\frac{H'(\underline{e})}{c'(\underline{e})}$, it follows that $L'_0(e)$ remains non-positive over the entire interval $[0, \underline{e}]$. This implies that $L_0(e)$ is decreasing on $[0, \underline{e}]$. Therefore, $e_0^* = \underline{e}$.
- (2) If $L'_0(0) < 0$ and $L'_0(\underline{e}) > 0$, that is $-\frac{H'(\underline{e})}{c'(\underline{e})} < \beta_0 < -\frac{H'(0)}{c'(0)}$, there must exist a unique $e_0^* \in (0, \underline{e})$ where $L'_0(e_0^*) = 0$. $L_0(e)$ is decreasing on $[0, e_0^*]$ and increasing on $[e_0^*, \underline{e}]$, implying that e_0^* is the unique minimum in this interval.
- (3) If $L'_0(0) \geq 0$, that is $\beta_0 \geq -\frac{H'(0)}{c'(0)}$, it follows that $L_0(e)$ is increasing on $[0, \underline{e}]$. Therefore, $e_0^* = 0$.

Thus, the lemma is proved. ■

When $a^*(e) = 0$, the insured retains the entire risk on his own. Then, according to Lemma 3.3, the pricing factor β_0 determines the demand for prevention. If effort is too costly, the individual will not exert effort. If the price is not excessive, some prevention is in demand (i.e. $e_0^* > 0$). The following lemma addresses the case $a^*(e) = 1$.

Lemma 3.4: *We have*

- In Case (i), we can obtain:

(1) If $\beta = 0$, then $e_2^* = +\infty$;

(2) If $0 < \beta < -\frac{(\theta_2\mu_0p(\bar{e})+\theta_1)\mu_0p'(\bar{e})}{c'(\bar{e})}$, then there uniquely exists $e_2^* \in (\bar{e}, +\infty)$, where e_2^* satisfies

$$(\theta_2\mu_0p(e_2^*) + \theta_1)\mu_0p'(e_2^*) + \beta c'(e_2^*) = 0;$$

(3) If $\beta \geq -\frac{(\theta_2\mu_0p(\bar{e})+\theta_1)\mu_0p'(\bar{e})}{c'(\bar{e})}$, then $e_2^* = \bar{e}$.

- In Case (ii), we can obtain:

(1) If $\beta = 0$, then $\tilde{e}_2^* = +\infty$;

(2) If $0 < \beta < -\frac{(\theta_2\mu_0p(\tilde{e})+\theta_1)\mu_0p'(\tilde{e})}{c'(\tilde{e})}$, then there uniquely exists $\tilde{e}_2^* \in (\tilde{e}, +\infty)$, where $\tilde{e}_2^*$ satisfies

$$(\theta_2\mu_0p(\tilde{e}_2^*) + \theta_1)\mu_0p'(\tilde{e}_2^*) + \beta c'(\tilde{e}_2^*) = 0;$$

(3) If $\beta \geq -\frac{(\theta_2\mu_0p(\tilde{e})+\theta_1)\mu_0p'(\tilde{e})}{c'(\tilde{e})}$, then $\tilde{e}_2^* = \tilde{e}$.

Proof: We know that

$$L_1(e) = \frac{\theta_2}{2}\mu_0^2p^2(e) + \theta_1\mu_0p(e) + \beta c(e).$$

Firstly, we consider Case (i). The analysis for Case (ii) follows similarly. We start by analyzing the derivatives of $L_1(e)$:

$$L_1'(e) = (\theta_2\mu_0p(e) + \theta_1)\mu_0p'(e) + \beta c'(e),$$

$$L_1''(e) = (\theta_2\mu_0p(e) + \theta_1)\mu_0p''(e) + \theta_2\mu_0^2p'(e)^2 + \beta c''(e) > 0.$$

This shows that $L_1'(e)$ is increasing on the interval $[\bar{e}, +\infty)$.

If $\beta = 0$, then $L_1'(e) \leq 0$ implying $e_2^* = +\infty$. If $\beta > 0$, given that $c''(e) > 0$, we have:

$$L_1'(+\infty) = +\infty > 0.$$

Now consider two cases:

- (1) If $L_1'(\bar{e}) \geq 0$, that is $\beta \geq -\frac{(\theta_2\mu_0p(\bar{e})+\theta_1)\mu_0p'(\bar{e})}{c'(\bar{e})}$, it follows that $L_1(e)$ is increasing over the entire interval $[\bar{e}, +\infty)$. This implies $e_1^* = \bar{e}$.
- (2) If $L_1'(\bar{e}) < 0$, that is $\beta < -\frac{(\theta_2\mu_0p(\bar{e})+\theta_1)\mu_0p'(\bar{e})}{c'(\bar{e})}$, then there must uniquely exist $e_1^* \in (\bar{e}, +\infty)$ such that $L_1'(e_1^*) = 0$. $L_1(e)$ is decreasing on $[\bar{e}, e_1^*]$ and increasing on $[e_1^*, +\infty)$, implying e_1^* is the unique minimum in this interval.

Thus, the lemma is proved. ■

Lemma 3.4 addresses the scenario where $a^*(e) = 1$. If $\beta = 0$, meaning that the insurer covers all effort costs, the insured will exert an infinite amount of effort. Otherwise, the optimal effort will be finite. It will either occur at the boundary point $\bar{e}$ (or $\tilde{e}$) or within the interval $(\bar{e}, +\infty)$ (or $(\tilde{e}, +\infty)$). Next, we analyze the case where $a^*(e) \in (0, 1)$.

Lemma 3.5: *Since $L_{a^*}(e)$ is continuous with respect to e , we can have*

- In Case (i), there exists e_1^* in the interval $[e, \bar{e}]$;

- In Case (ii), there exists $\tilde{e}_1^*$ in the interval $[0, \bar{e}]$.

Since $L_{a^*}(e)$ is continuous on $[\underline{e}, \bar{e}]$, Lemma 3.5 follows directly. Now we identify the globally optimal effort e^* as follows.

Theorem 3.6: Under Assumption 3.1, we summarize the globally optimal effort e^* as follows:

- (1) In Case (i), we have:

$$e^* = \arg \min_{e \in \{e_0^*, e_1^*, e_2^*\}} \{L_0(e_0^*), L_{a^*}(e_1^*), L_1(e_2^*)\}.$$

Here, e_0^* , e_1^* , and e_2^* are defined in Lemmas 3.3, 3.5 and 3.4.

- (2) In Case (ii), we have:

$$e^* = \arg \min_{e \in \{\tilde{e}_1^*, \tilde{e}_2^*\}} \{L_{a^*}(\tilde{e}_1^*), L_1(\tilde{e}_2^*)\}.$$

In this case, $\tilde{e}_1^*$ and $\tilde{e}_2^*$ are defined in Lemmas 3.5 and 3.4.

For Case (i), it is straightforward to verify that the scenarios where $e_1^* = \underline{e}$ or $e_1^* = \bar{e}$ are equivalent to $e_0^* = \underline{e}$ or $e_2^* = \bar{e}$, respectively. This equivalence arises because $L_0(e)$ and $L_1(e)$ coincide with $L_{a^*}(e)$ on the intervals $[0, \underline{e}]$ and $[\bar{e}, +\infty)$, respectively, leading to the identities $L_{a^*}(\underline{e}) = L_0(\underline{e})$ and $L_{a^*}(\bar{e}) = L_1(\bar{e})$.

However, the objective function is not differentiable at $\underline{e}$ and $\bar{e}$. This is because the form of $a^*(e)$ changes at these points, making the left and right derivatives different. Moreover, since $L_{a^*}(e)$ may not be convex,⁶ we cannot establish sufficient conditions for the existence of unique solution as in Lemmas 3.3 and 3.4. Then, the uniqueness of the global optimal solution cannot be guaranteed. Although we cannot guarantee that e_1^* or e^* is unique, the proposition below gives conditions for when e^* equals e_1^* and shows how to compute e_1^* .

Proposition 3.7: Under Assumption 3.1, e^* can be determined as follows:

- In Case (i), if $\beta_0 \leq -\frac{H'(\underline{e})}{c'(\underline{e})}$ and $\beta \geq -\frac{(\theta_2 \mu_0 p(\bar{e}) + \theta_1) \mu_0 p'(\bar{e})}{c'(\bar{e})}$, then $e^* = e_1^*$.
- In Case (ii), if $\beta \geq -\frac{(\theta_2 \mu_0 p(\bar{e}) + \theta_1) \mu_0 p'(\bar{e})}{c'(\bar{e})}$, then $e^* = \tilde{e}_1^*$.

In both cases, e^* is the solution to the equation:

$$\begin{aligned} & a^*(e)^2 \left[\theta_2 p(e) p'(e) \mu_0^2 + \frac{\gamma_1}{2} p'(e) (\lambda_0 - 2p(e) \mu_0^2) \right] \\ & + a^*(e) [(\theta_1 - 1) p'(e) \mu_0 - \gamma_1 p'(e) (\lambda_0 - 2p(e) \mu_0^2) - (\beta_0 - \beta) c'(e)] \\ & + p'(e) \left(\mu_0 + \frac{\gamma_1}{2} \lambda_0 - p(e) \gamma_1 \mu_0^2 \right) + \beta_0 c'(e) = 0, \end{aligned} \quad (15)$$

where $a^*(e)$ is given by Equation (12).

Proof: We first consider Case (i). The proof for Case (ii) is similar. In Case (i), the proof has two steps. First, we give the condition under which $e^* = e_1^*$. Second, we provide a semi-closed form of e^* .

By Lemma 3.3, if $L'_0(\underline{e}) \leq 0$, then $e^* \geq \underline{e}$. By Lemma 3.4, if $L'_1(\bar{e}) \geq 0$, then $e^* \leq \bar{e}$. Hence, if both $L'_0(\underline{e}) \leq 0$ and $L'_1(\bar{e}) \geq 0$ hold, we obtain $e^* \in [\underline{e}, \bar{e}]$. In this case, the optimal $a^*(e) \in [0, 1]$ and is given by (12).

⁶ This is further substantiated by the numerical results in Figure 6, see Section 5.

Next, we derive the semi-closed form of e^* . The key idea is that, by optimality, e^* must satisfy both the expression for $a^*(e)$ and that of $e^*(a)$. In other words, it lies at the intersection of $e^*(a)$ and $a^*(e)$. Since $e^* \in [\underline{e}, \bar{e}]$, at least one such intersection must exist.

Now we give the expression for $e^*(a)$. To proceed, we calculate the derivative of the loss function $L_a(e)$ with respect to e for a given a . We can obtain:

$$\begin{aligned} L'_a(e) &= (1 - a + \theta_1 a) \mathbb{E}'[X_e] + \theta_2 a^2 \mathbb{E}[X_e] \mathbb{E}'[X_e] + [\beta_0 - a(\beta_0 - \beta)] c'(e) + \frac{\gamma_1}{2} (1 - a)^2 \text{Var}'[X_e], \\ L''_a(e) &= (1 - a) \left[\mathbb{E}''[X_e] + \frac{\gamma_1}{2} (1 - a) \text{Var}''[X_e] \right] + \theta_2 a^2 (\mathbb{E}'[X_e])^2 + (\theta_2 a \mathbb{E}[X_e] + \theta_1 a) \mathbb{E}''[X_e] \\ &\quad + [(1 - a)\beta_0 + a\beta] c''(e). \end{aligned}$$

In the following we prove $L''_a(e) > 0$. Since $\mathbb{E}''[X_e] > 0$, if $\text{Var}''[X_e] > 0$, we have $\mathbb{E}''[X_e] + \frac{\gamma_1}{2} (1 - a) \text{Var}''[X_e] > 0$. If $\text{Var}''[X_e] < 0$, assumption $H''(e) > 0$ ensures that the first term of $L''_a(e)$ is positive. Thus, $L''_a(e) > 0$, meaning that for any given a , $L'_a(e)$ is an increasing function of e .

Now, we examine the behavior of $L'_a(e)$ at the boundaries:

(1) At $e = 0$, we have:

$$L'_a(0) = (1 - a) \left[\mathbb{E}'[X_0] + \frac{\gamma_1}{2} (1 - a) \text{Var}'[X_0] \right] + (\theta_2 a^2 \bar{p} \mu_0 + \theta_1 a) \mathbb{E}'[X_0].$$

Since $\mathbb{E}'[X_0] < 0$ and $H'(0) = \mathbb{E}'[X_0] + \frac{\gamma_1}{2} \text{Var}'[X_0] < 0$, we know that $L'_a(0) < 0$.

(2) At $e \rightarrow +\infty$, we have:

$$L'_a(+\infty) = [\beta_0 - a(\beta_0 - \beta)] c'(+\infty) > 0.$$

Thus, for any given a , there uniquely exists $e^*(a) \in (0, +\infty)$ such that $L'_a(e^*(a)) = 0$, and $L_a(e)$ is decreasing in $[0, e^*(a)]$ and increasing in $[e^*(a), +\infty)$.

Now, by combining $a^*(e)$ and $e^*(a)$, we know that the optimal e^* satisfies $L'_{a^*(e)}(e) = 0$, which can be rewritten as:

$$\begin{aligned} a^*(e)^2 \left[\theta_2 \mathbb{E}[X_e] \mathbb{E}'[X_e] + \frac{\gamma_1}{2} \text{Var}'[X_e] \right] + a^*(e) \left[(\theta_1 - 1) \mathbb{E}'[X_e] - \gamma_1 \text{Var}'[X_e] \right] \\ - (\beta_0 - \beta) c'(e) + H'(e) + \beta_0 c'(e) = 0, \end{aligned} \quad (16)$$

where $a^*(e)$ is given by Equation (12). Substituting the detailed expressions for $\mathbb{E}[X_e]$ and $\text{Var}[X_e]$ into this equation leads to the desired condition (15).

For Case (ii), the proof follows the same reasoning. ■

In the following section, we will address the insurer's problem (2.2) using the pair $(a^*(e^*), e^*)$ derived in the current section.

4. Solution to the insurer's problem

In this section, we further analyze the insurer's problem based on the previously obtained pair $(a^*(e^*), e^*)$, which depends on the value of β . For convenience, we denote this pair as $(a^*(\beta), e^*(\beta))$. Accordingly, the thresholds that partition the domain of e are also functions of β , that is, $\underline{e}(\beta)$, $\bar{e}(\beta)$, $\tilde{e}(\beta)$. Then, the insurer's problem (2.2) is equivalent to solving the following optimization

problem:

$$\sup_{\beta \in [0,1]} J(\beta), \quad (17)$$

where

$$J(\beta) = \frac{\theta_2 a^*(\beta)^2}{2} \mathbb{E}[X_{e^*(\beta)}]^2 + (\theta_1 - 1) a^*(\beta) \mathbb{E}[X_{e^*(\beta)}] - a^*(\beta)(1 - \beta)c(e^*(\beta)).$$

Before stating the solution of problem (17), we first discuss the uniqueness of $(a^*(\beta), e^*(\beta))$. In the previous section, we have shown only the existence of $(a^*(\beta), e^*(\beta))$, but not its uniqueness. Since the objective function may not be convex in the decision variables, the uniqueness of the optimal pair cannot be guaranteed in general. However, in practice, even if Problem 2.1 admits multiple solutions, one must still choose a specific solution when determining β^* .

After selecting $(a^*(\beta), e^*(\beta))$, the range of $e^*(\beta)$ partitions the domain of β into multiple intervals. Due to the mathematical complexity of $e^*(\beta)$, explicit expressions for these intervals are not available. Instead, the intervals can only be described implicitly using inequalities that depend on β . This leads to the following two lemmas:

Lemma 4.1: Consider the objective function $J(\beta)$ in (17), in Case (i), if $0 \leq e^*(\beta) \leq \underline{e}(\beta)$, then $J(\beta)$ does not depend on β .

Proof: In Case (i), if $0 \leq e^*(\beta) \leq \underline{e}(\beta)$, then $a^*(\beta) = 0$ and the objective function satisfies $J(\beta) = 0$, which is independent of β . The same applies to Case (ii). ■

Lemma 4.2: Consider the objective function $J(\beta)$ in (17). If either of the following two conditions holds:

- In Case (i), $\bar{e}(\beta) < e^*(\beta) < +\infty$; or
- In Case (ii), $\bar{e}(\beta) < e^*(\beta) < +\infty$,

then $J(\beta)$ is increasing in β .

Proof: We begin by discussing Case (i). When $\bar{e}(\beta) < e^*(\beta) < +\infty$, we have $a^*(\beta) = 1$, and $e^*(\beta)$ satisfies

$$L'_1(e^*(\beta)) = (\theta_2 \mu_0 p(e^*(\beta)) + \theta_1) \mu_0 p'(e^*(\beta)) + \beta c'(e^*(\beta)) = 0. \quad (18)$$

The objective function takes the form:

$$J(\beta) = \frac{\theta_2}{2} p(e^*(\beta))^2 \mu_0^2 + (\theta_1 - 1) p(e^*(\beta)) \mu_0 - (1 - \beta) c(e^*(\beta)).$$

We compute the derivative of $J(\beta)$ with respect to β :

$$\frac{dJ(\beta)}{d\beta} = \frac{de^*(\beta)}{d\beta} [\theta_2 \mu_0^2 p(e^*(\beta)) p'(e^*(\beta)) + (\theta_1 - 1) \mu_0 p'(e^*(\beta)) - (1 - \beta) c'(e^*(\beta))] + c(e^*(\beta)).$$

By (18), we can obtain:

$$\frac{dL(\beta)}{d\beta} = -\frac{de^*(\beta)}{d\beta} [\mu_0 p'(e^*(\beta)) + c'(e^*(\beta))] + c(e^*(\beta)).$$

Since $\beta \leq 1 < \theta_1 < \theta_2 \mu_0 p(e^*(\beta)) + \theta_1$, we find that:

$$\mu_0 p'(e^*(\beta)) + c'(e^*(\beta)) = \mu_0 p'(e^*(\beta)) \left(1 - \frac{\theta_2 \mu_0 p(e^*(\beta)) + \theta_1}{\beta} \right) > 0.$$

Next, we differentiate both sides of Equation (18) with respect to β :

$$\frac{de^*(\beta)}{d\beta} = -\frac{c'(e^*(\beta))}{\theta_2\mu_0^2p'(e^*(\beta))^2 + (\theta_2\mu_0p(e^*(\beta)) + \theta_1)\mu_0p''(e^*(\beta)) + \beta c''(e^*(\beta))} < 0.$$

Thus, we have $\frac{dL(\beta)}{d\beta} > 0$. Then, *Case (ii)* can be checked similarly. ■

Lemma 4.1 addresses the case where $a^*(\beta) \equiv 0$, and Lemma 4.2 considers the case where $a^*(\beta) \equiv 1$. Lemma 4.2 shows that when the insured buys full coverage, the insurer's optimal strategy is to minimize the shared self-protection costs. This conclusion is intuitive, as raising β in this case does not raise premium income but only increases the insurer's costs. Conversely, Lemma 4.1 indicates that if the insured never buys insurance, β does not affect the insurer's wealth.

Based on Lemmas 4.1 and 4.2, we can characterize β^* in some cases.

Theorem 4.3: *Consider the insurer's Problem 2.2, we have*

- In *Case (i)*, for all $\beta \in [0, 1]$,
 - (1) If $0 \leq e^*(\beta) \leq \underline{e}(\beta)$, then the insurer is indifferent over the value of β ;
 - (2) If $\underline{e}(\beta) \leq e^*(\beta) \leq \bar{e}(\beta)$, there exists β^* in $[0, 1]$;
 - (3) If $\bar{e}(\beta) < e^*(\beta) < +\infty$, then $\beta^* = 1$.
- In *Case (ii)*, for all $\beta \in [0, 1]$,
 - (1) If $0 \leq e^*(\beta) \leq \tilde{e}(\beta)$, there exists $\beta^* \in [0, 1]$;
 - (2) If $\tilde{e}(\beta) < e^*(\beta) < +\infty$, then $\beta^* = 1$.

Proof: We begin with *Case (i)*. Results (1) and (3) follow directly from Lemmas 4.1 and 4.2. Result (2) is also immediate from the continuity of $J(\beta)$ in β . *Case (ii)* can be verified similarly. ■

The cases in Theorem 4.3 are highly specific. Therefore, we will further investigate the properties of β^* through numerical analysis in Section 5.

5. Numerical analysis

In this section, we will employ numerical analysis to present the theoretical results and demonstrate conclusions that are challenging to establish theoretically. We will first introduce the parameter settings for the analysis, followed by an examination of the insured's problem and the insurer's problem. For the numerical analysis design, we suppose the random variable Y to follow a Gamma distribution $\Gamma(k, \theta)$. The exponentially decaying function $p(e)$ is defined as:

$$p(x) = (\bar{p} - \underline{p})e^{-x} + \underline{p},$$

To ensure that Assumption 3.1 holds, we set the parameters as follows: $k = 0.35$, $\theta = 2.5$, $\bar{p} = 0.99$, and $\underline{p} = 0.01$. We adopt a quadratic cost function defined by:

$$c(e) = \frac{e^2}{5}.$$

The other baseline parameters are set to: $\theta_1 = 6.5$, $\theta_2 = 9$, $\gamma_1 = 2$, $\beta_0 = 1.5$, and $\beta = 0.8$. The parameter selection in this paper is not based on empirical data. Instead, it aims to represent diverse scenarios within a single figure.

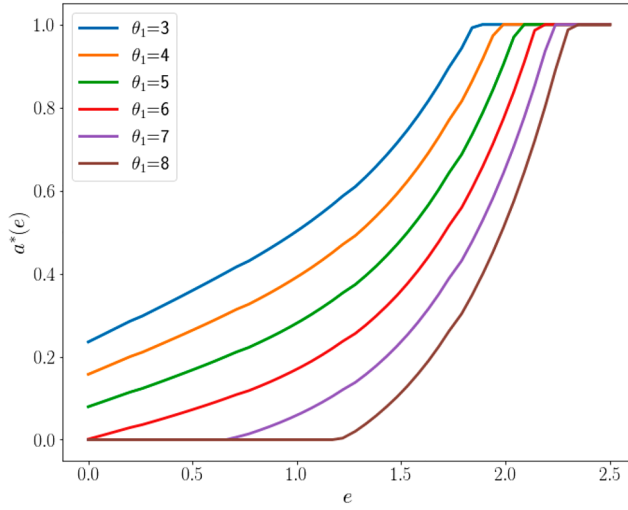


Figure 3. Monotonicity of $a^*(e)$ with respect to e .

5.1. Monotonicity of $a^*(e)$ with respect to e

Firstly, we investigate the conclusions in Proposition 3.2, particularly the monotonicity of $a^*(e)$ with respect to e . Although Figure 2 illustrates the shape of $a^*(e)$, it is only a schematic, and we have not theoretically proven the monotonicity of $a^*(e)$ with respect to e when $a^*(e) \in (0, 1)$. Based on the parameter values provided earlier, we know that:

$$\gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right) = 5.0175.$$

We now examine the shape of $a^*(e)$ by setting $\theta_1 = [3, 4, 5, 6, 7, 8]$. The results are shown in Figure 3, where the horizontal axis represents e and the vertical axis represents the values of $a^*(e)$. Each of the six curves in the figure corresponds to a different value of θ_1 . For $\theta_1 = [3, 4, 5]$, we have:

$$\theta_1 - 1 < \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right).$$

According to Lemma 3.1, only $\bar{e}$ exists in this case. For $\theta_1 = [6, 7, 8]$, we have:

$$\theta_1 - 1 > \gamma_1 \left(\frac{\lambda_0}{\mu_0} - \bar{p}\mu_0 \right),$$

where $\underline{e}$ and $\bar{e}$ exist.

Figure 3 shows that, in both cases, $a^*(e)$ is increasing in e . Furthermore, for small values of e , $a^*(e)$ exhibits concavity, while as $a^*(e)$ approaches 1, it shows convexity. Overall, $a^*(e)$ displays an S shape with respect to e .

This implies that, given a certain level of self-protection, the optimal insurance demand of the insured increases with effort. The marginal effect of effort on insurance is first increasing, then decreasing.

5.2. Effect of θ_1 and θ_2 on e^* and $a^*(e^*)$

Now, we analyze the impact of the premium loading factors θ_1 and θ_2 on the optimal self-protection effort level e^* and the optimal insurance demand $a^*(e^*)$. The results are shown in Figure 4, which

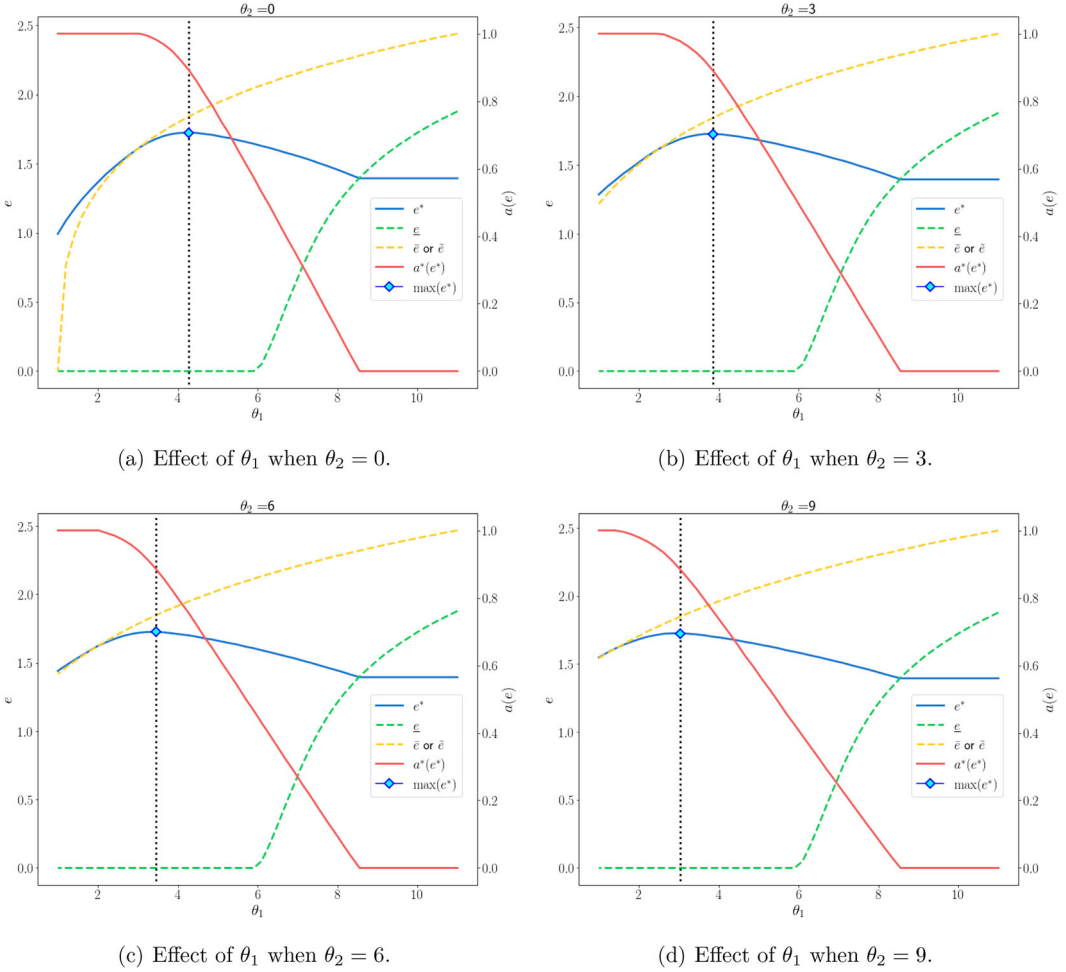


Figure 4. Effect of premium loading on insurance demand and optimal effort. (a) Effect of θ_1 when $\theta_2 = 0$. (b) Effect of θ_1 when $\theta_2 = 3$. (c) Effect of θ_1 when $\theta_2 = 6$ and (d) Effect of θ_1 when $\theta_2 = 9$.

contains four subplots corresponding to different values of $\theta_2 = [0, 3, 6, 9]$. In each subplot, the horizontal axis represents θ_1 , with a range from 1 to 11. The left vertical axis corresponds to the value of e , while the right vertical axis represents $a(e)$.

In the figure, the solid blue line illustrates how e^* varies with θ_1 , and we mark the maximum value of e^* with a blue diamond. The green dashed line represents $\underline{e}$, and the yellow dashed line represents $\bar{e}$ or $\tilde{e}$. Notice that as θ_1 increases, we first encounter the scenario where only $\bar{e}$ exists, followed by the scenario where both $\underline{e}$ and $\bar{e}$ exist. Therefore, for simplicity, when only $\bar{e}$ exists, we set $\underline{e} = 0$, which does not contradict our theoretical results. Additionally, since $\bar{e}$ and $\tilde{e}$ are determined by the same equation, representing them with the same dashed line is reasonable. We use a solid red line to indicate $a^*(e^*)$.

Next, by examining each subplot individually, we observe that e^* initially increases with θ_1 , then decreases, and eventually becomes constant. This indicates that as the premium becomes more expensive, the optimal level of self-protection initially rises, then falls, and finally stabilizes. The last case is straightforward to understand. From the figures, we see that when $e^* > \underline{e}$, this final situation occurs. At this point, $a^*(e^*) = 0$, which removes β from the objective function $L(a, e)$, so e^* no longer changes.

Moreover, by comparing e^* with the green and yellow dashed lines, we observe that when e^* exceeds $\bar{e}$ or $\tilde{e}$, $a^*(e^*) = 1$, and when e^* is below $\underline{e}$, $a^*(e^*) = 0$. This is consistent with the conclusions in Proposition 3.2. The case where $a^*(e^*) \in (0, 1)$ is more complex. We find that to the left of the black dashed line, where e^* is increasing, $a^*(e^*)$ decreases. This suggests that as premiums become more expensive, the optimal level of self-protection increases while the optimal insurance demand decreases, indicating a substitution effect between self-protection and insurance when premiums are relatively affordable.

However, to the right of the black dashed line, where e^* begins to decrease, $a^*(e^*)$ continues to decline. This indicates that as premiums become even more expensive, both the optimal self-protection level and optimal insurance demand decrease together, suggesting a complementary effect between self-protection and insurance when premiums are high. Under the cost-sharing mechanism, our results extend the classical findings in Ehrlich and Becker (1972), showing that self-protection and market insurance can act as either substitutes or complements even when the effort is fully reflected in the premium.

Comparing the four subplots, two trends emerge as θ_2 increases. First, $a^*(e^*)$ decreases, which is intuitive since insurance demand declines as premiums rise. Second, the peak of e^* shifts to the left as θ_2 increases, indicating that as premiums become more costly, the complementary effect between self-protection and insurance is more likely to occur. This conclusion aligns with the insights from observing each subplot individually.

5.3. Effect of β on e^* and $a^*(e^*)$

In this section, we analyze the effect of the cost-sharing ratio, β , on e^* and $a^*(e^*)$. The results of the numerical analysis are shown in Figure 5, which comprises four subplots, each corresponding to a different value of θ_1 , specifically $\theta_1 = [3, 5, 7, 9]$. In each subplot, the horizontal axis represents the cost-sharing ratio β , ranging from 0 to 1, while the left vertical axis shows e^* , and the right vertical axis represents $a^*(e^*)$. The curve types, colors, and corresponding variables in the legend match those in Figure 4.

Starting with the subplot Figure 5(a), we observe that both $a^*(e^*)$ and e^* decrease as β increases. This suggests that as the insured's share of the cost rises, their demand for insurance and willingness to engage in self-protection both decline. This result is intuitive: on one hand, an increased cost burden for a given effort level e^* discourages the insured party from engaging in self-protection due to the higher associated expenses. On the other hand, cost-sharing by the insurer serves as an incentive for the insured party to increase their demand for insurance. Therefore, as the insurer's share of the cost decreases, the insured party's demand for insurance diminishes accordingly. Notably, the joint decline in both $a^*(e^*)$ and e^* indicates a complementary relationship between insurance demand and self-protection in this context. While the observation that effort decreases as the cost burden rises is consistent with the law of demand, prevention can be a Giffen good in extreme cases (see Peter, 2021b).

Examining subplots 5(b–d), we find that the previously noted relationship between $a^*(e^*)$ and e^* still holds. Additionally, we observe that the changes in e^* with respect to β are not continuous, leading to jumps in $a^*(e^*)$ as β varies. Specifically, these jumps occur when the curve representing e^* crosses either the $\bar{e}$ or $\tilde{e}$ and the $\underline{e}$ lines. This phenomenon arises because the objective function, $L(a^*(e), e)$, is not convex over the entire range of e . Given that the domain of e is divided into two or three intervals, the optimal e^* corresponds to the local optimum within one of these intervals. As β changes, e^* may shift between different intervals.

An example of this discontinuity is illustrated in Figure 6, which corresponds to the scenario in subplot Figure 5(d). In this example, as β increases from 0.61 to 0.62, e^* jumps from a value above $\bar{e}$ to below $\underline{e}$. Subplot Figure 6(a) reveals that when $\beta = 0.61$, e^* is located within the third interval, $[\bar{e}, +\infty)$. However, as β increases to 0.62, e^* shifts to the first interval, $[0, \underline{e}]$. Since $L(a^*(e), e)$

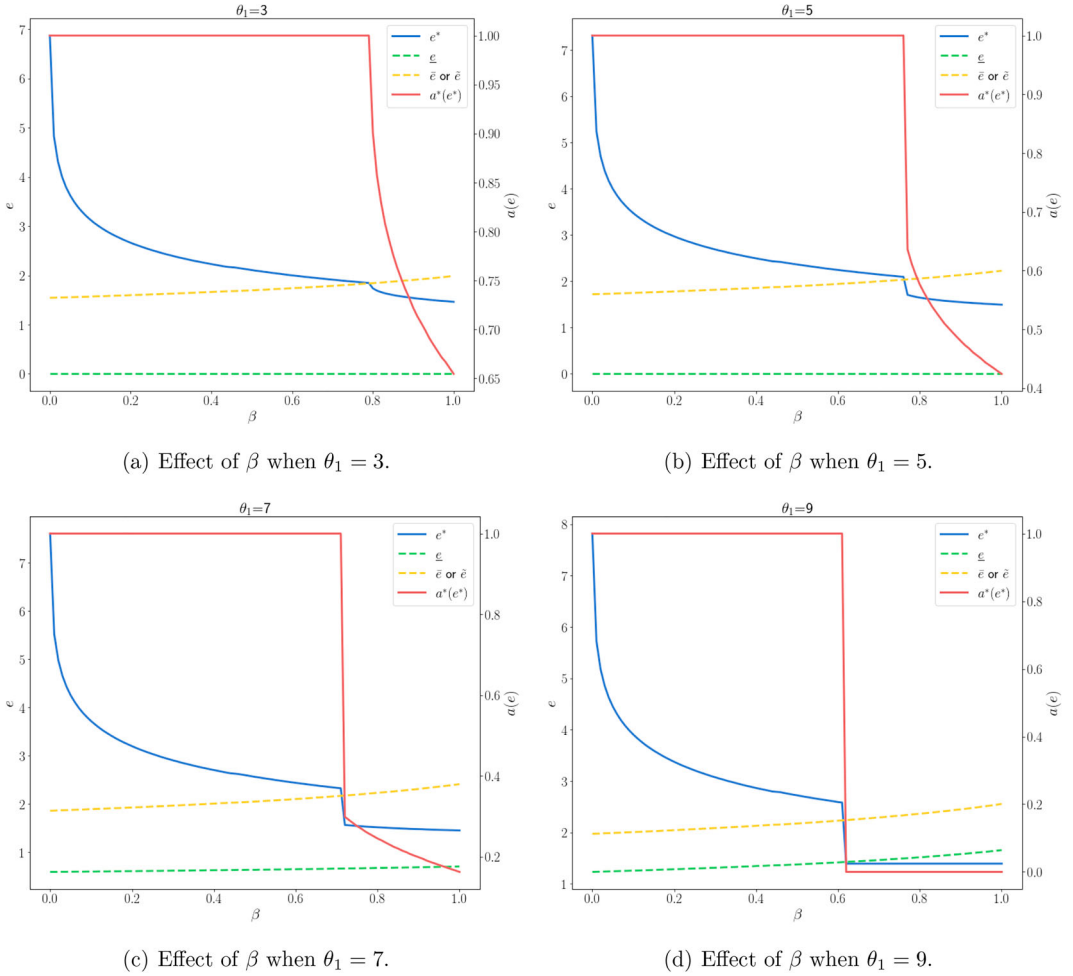


Figure 5. Effect of cost-sharing proportion on insurance demand and optimal effort. (a) Effect of β when $\theta_1 = 3$. (b) Effect of β when $\theta_1 = 5$. (c) Effect of β when $\theta_1 = 7$ and (d) Effect of β when $\theta_1 = 9$.

is non-convex, e^* will jump between local optima across intervals at certain critical points, moving from one interval's local optimum to that of another.

5.4. Effect of γ_1 on e^* and $a^*(e^*)$

In this section, we examine the impact of the risk aversion parameter γ_1 on the optimal self-protection level e^* and insurance demand $a^*(e^*)$. The results are illustrated in Figure 7, which includes four subplots representing different values of θ_1 , specifically $\theta_1 = [3, 5, 7, 9]$. Each subplot shows γ_1 on the x-axis, ranging from 0 to 4, with e^* on the left y-axis and $a^*(e^*)$ on the right y-axis. The legend follows the same format as in Figure 4.

From Figure 7, we observe that both e^* and $a^*(e^*)$ increase as γ_1 grows, indicating that higher risk aversion leads to greater insurance demand and stronger self-protection. This result aligns intuitively with expectations and shows a complementary relationship between insurance demand and self-protection in this context. Additionally, in standard expected utility models, the monotonic relationship between risk aversion and effort is ambiguous due to opposing forces from risk aversion and downside risk aversion; see Peter (2025) for more studies.

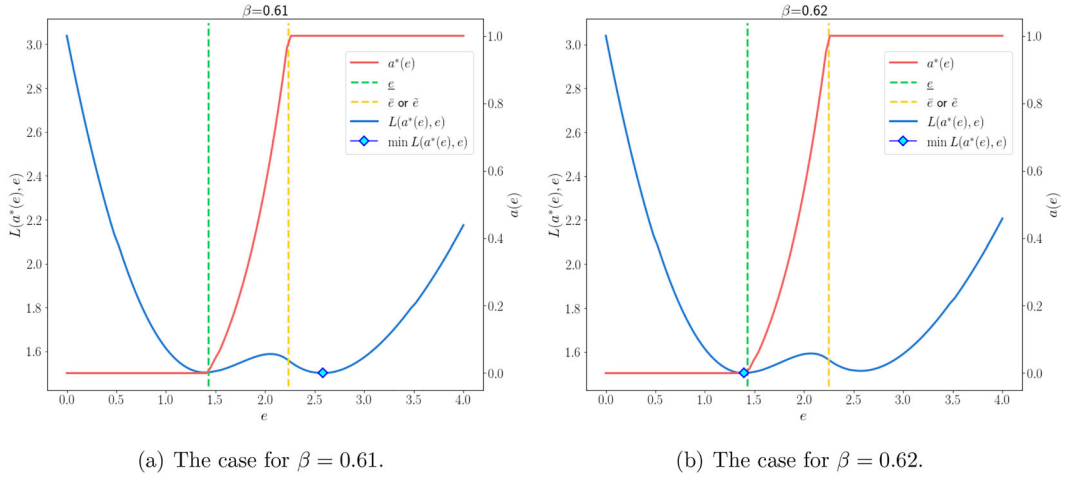


Figure 6. Explanation for the jump of e^* . (a) The case for $\beta = 0.61$ and (b) The case for $\beta = 0.62$.

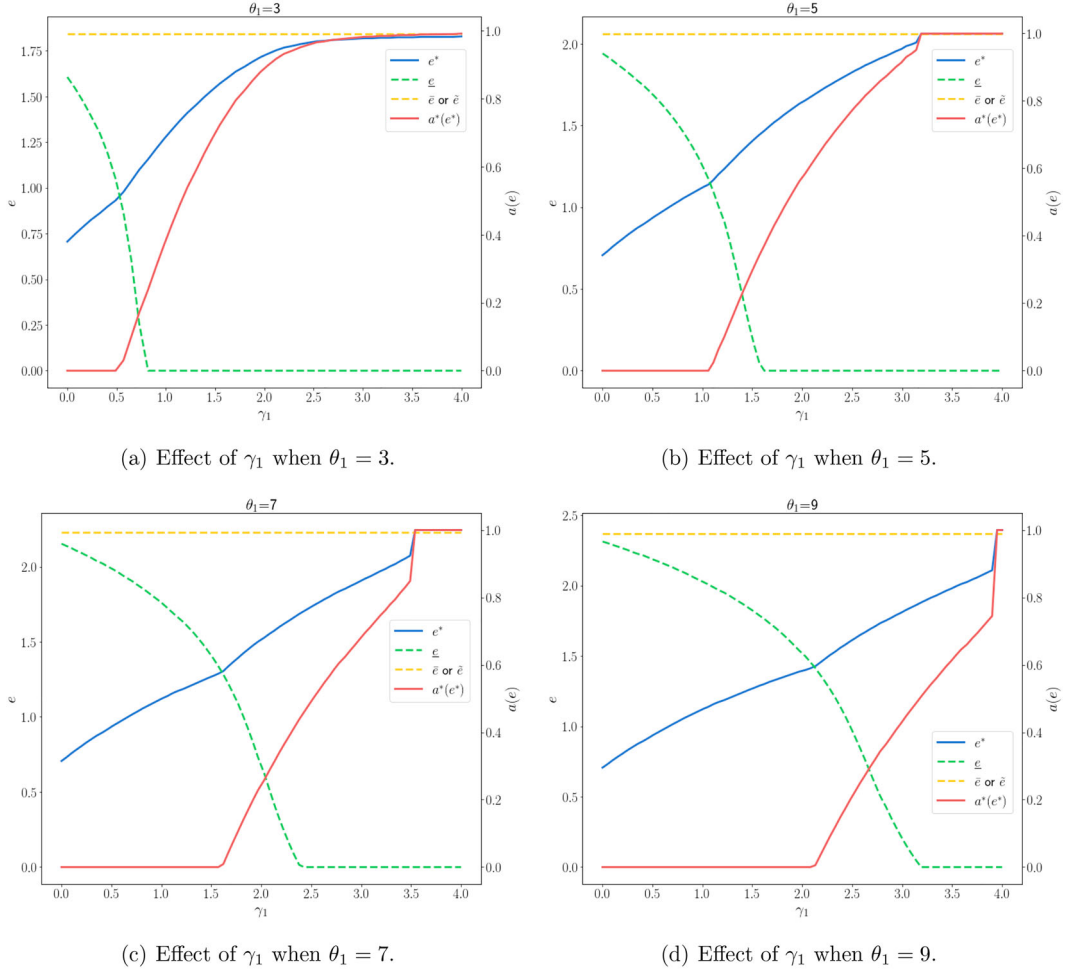


Figure 7. Effect of risk aversion on insurance demand and optimal effort. (a) Effect of γ_1 when $\theta_1 = 3$. (b) Effect of γ_1 when $\theta_1 = 5$. (c) Effect of γ_1 when $\theta_1 = 7$ and (d) Effect of γ_1 when $\theta_1 = 9$.

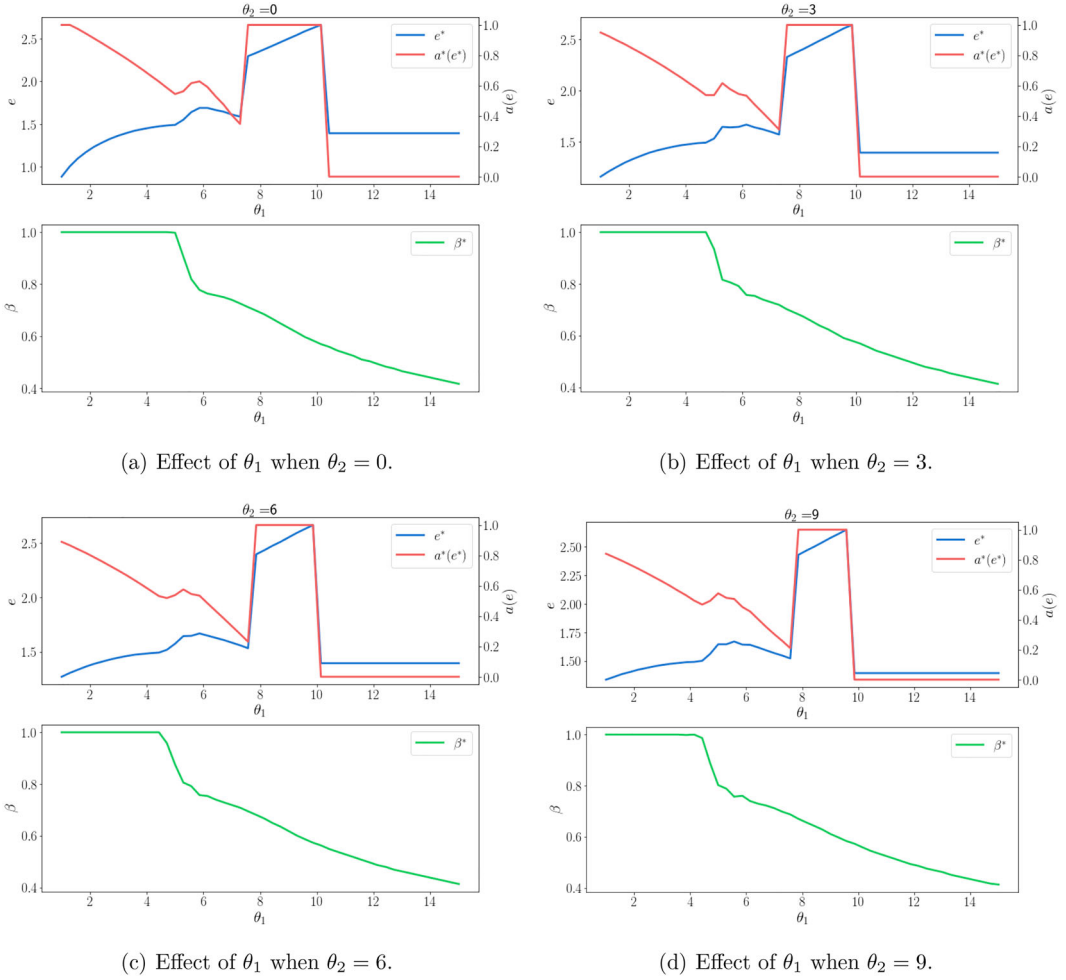


Figure 8. Effect of premium loading on cost-sharing ratio. (a) Effect of θ_1 when $\theta_2 = 0$. (b) Effect of θ_1 when $\theta_2 = 3$. (c) Effect of θ_1 when $\theta_2 = 6$ and (d) Effect of θ_1 when $\theta_2 = 9$.

As in Figure 5, both e^* and $a^*(e^*)$ exhibit discontinuities, specifically when the curve for e^* intersects with the $\bar{e}$ or $\tilde{e}$ lines. The underlying reason for these jumps mirrors that in Figure 5: non-convexities in the objective function across different intervals. We also notice that when the e^* curve intersects the $\bar{e}$ line, the changes in e^* and $a^*(e^*)$ are continuous but not smooth, with a sudden shift in growth rates on either side of $\bar{e}$. This phenomenon arises because the characteristics of the objective function change around $\bar{e}$. Specifically, on the left of $\bar{e}$, $a^*(e^*)$ is always zero, while on the right, $a^*(e^*)$ is a function of both e^* and γ_1 . For further details, refer to Proposition 3.2.

5.5. Effect of θ_1 and θ_2 on β^*

This section analyzes the optimal decision of the insurer, particularly the effects of premium loading parameters θ_1 and θ_2 on the optimal cost-sharing ratio, β^* . The results are illustrated in Figure 8, which comprises four subplots corresponding to different values of $\theta_2 = [0, 3, 6, 9]$. Each subplot consists of two panels, providing insights into the behavior of optimal insurance demand $a^*(e^*)$, the optimal self-protection effort e^* , and the optimal cost-sharing ratio β^* .

The horizontal axis in each subplot represents θ_1 , ranging from 0 to 15. In the upper panels, the left vertical axis measures e^* at $\beta = \beta^*$, while the right vertical axis measures $a^*(e^*)$ under the same condition. The lower panels depict the optimal cost-sharing ratio, β^* , on their vertical axis. This layout provides a comprehensive visualization of how changes in premium loading parameters influence the insurer's decisions and the resulting behavior of the insured.

As θ_1 increases, the optimal cost-sharing ratio β^* decreases monotonically. This trend indicates that as the premium rate rises, the insurer opts to bear a larger proportion of the self-protection costs, thereby reducing the cost burden on the policyholder. As shown in Section 5.2, an increase in θ_1 suppresses policyholders' insurance demand, leading to a decline in the insurer's premium income. To counteract this effect, and as discussed in Section 5.3, the insurer reduces β^* to incentivize policyholders to purchase more insurance, thereby increasing premium revenue. Consequently, the optimal cost-sharing ratio β^* decreases as θ_1 increases.

A comparison of the four subplots reveals an additional pattern: for the same value of θ_1 , higher values of θ_2 correspond to smaller (or at least not larger) values of β^* . This suggests that the second premium loading parameter, θ_2 , further reinforces the insurer's incentive to assume a greater share of self-protection costs as premium loadings rise.

However, when focusing on the upper panels of the subplots, it becomes evident that neither $a^*(e^*)$ nor e^* exhibits a monotonic relationship with θ_1 at $\beta = \beta^*$. This complexity highlights the nuanced interplay between premium parameters and policyholder behavior, suggesting that additional factors may influence insurance demand and self-protection decisions under varying cost-sharing conditions.

5.6. Effect of γ_1 on β^*

In this section, we analyze the impact of the policyholder's risk aversion coefficient, γ_1 , on the optimal cost-sharing ratio, β^* . While γ_1 does not explicitly appear in the insurer's objective function, it indirectly influences β^* through its effect on the policyholder's decisions regarding optimal coverage level, $a^*(e^*)$, and self-protection effort, e^* . These interactions play a crucial role in shaping the insurer's cost-sharing strategy. The results of this analysis are presented in Figure 9.

Figure 9 consists of four subplots, each corresponding to a different value of θ_1 , specifically $\theta_1 = [3, 5, 7, 9]$. Within each subplot, two panels provide a comprehensive depiction of the results. The upper panel illustrates the variation in $a^*(e^*)$ and e^* as γ_1 changes, highlighting how risk aversion influences both the policyholder's insurance demand and self-protection effort. The lower panel presents the corresponding changes in the optimal cost-sharing ratio, β^* , offering insights into how the insurer adjusts its cost-sharing strategy in response to variations in policyholder behavior. The horizontal axis in all subplots represents the range of γ_1 , which spans from 0 to 4. The structure and legend of Figure 9 are consistent with those of Figure 8, ensuring ease of comparison and interpretation.

By examining Figure 9, we observe a clear monotonic relationship between β^* and γ_1 , where β^* increases as γ_1 rises. This indicates that as the policyholder's level of risk aversion increases, the insurer assigns a larger share of self-protection costs to the policyholder. This outcome aligns with the results discussed in Section 5.4, which demonstrate that both $a^*(e^*)$ and e^* increase with higher levels of γ_1 . From the insurer's perspective, this adjustment in β^* serves to stimulate policyholders' demand for insurance, thereby increasing premium revenue. Although higher levels of self-protection effort (e^*) reduce premium revenue and increase the cost burden for policyholders, the corresponding increase in premium revenue driven by the rise in β^* consistently outweighs the revenue loss caused by increased e^* . As a result, the insurer's optimal strategy involves raising β^* in response to greater policyholder risk aversion (γ_1).

Finally, regarding the behavior of $a^*(e^*)$ and e^* when $\beta = \beta^*$, the findings are consistent with those presented in Section 5.5. Specifically, $a^*(e^*)$ and e^* do not exhibit a monotonic relationship with γ_1 . This suggests that while γ_1 exerts a clear influence on certain parameters, its effect on the

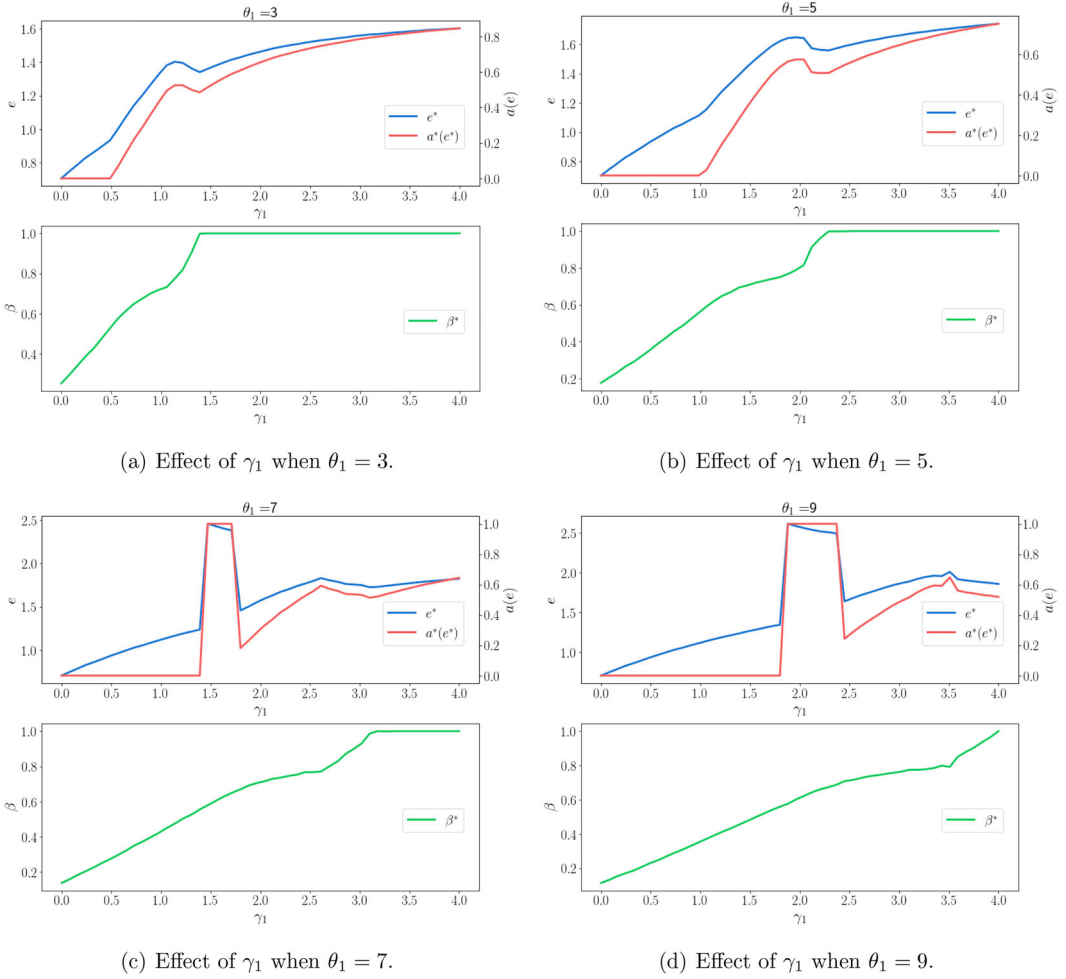


Figure 9. Effect of premium loading on cost-sharing ratio. (a) Effect of γ_1 when $\theta_1 = 3$. (b) Effect of γ_1 when $\theta_1 = 5$. (c) Effect of γ_1 when $\theta_1 = 7$ and (d) Effect of γ_1 when $\theta_1 = 9$.

optimal levels of $a^*(e^*)$ and e^* is more complex and warrants further exploration to fully capture the nuances of this interaction.

6. Conclusion

This paper explores a mean-variance Stackelberg game between an insured and an insurer, offering a comprehensive framework in which the insured can exert effort to reduce the probability of loss occurring, thereby lowering the insurance premium. The model incorporates proportional insurance under a quadratic convex premium principle and features a cost-sharing mechanism for self-protection expenses. In this game, the insurer is leader who strategically chooses the optimal cost-sharing proportion to maximize its expected wealth. The insured is follower who determines the optimal levels of insurance coverage and prevention effort, aiming to maximize a mean-variance functional of their wealth. We derived the explicit expressions for the optimal insurance proportion and protection level.

Furthermore, we conduct numerical analysis to explore the monotonic relationships between these variables and key model parameters. The analysis reveal how self-protection and insurance demand

may act as either complements or substitutes depending on the specific conditions, generalizing the existing results established in Ehrlich and Becker (1972), Bensalem et al. (2020), S. Chen et al. (2025) and Q. Li et al. (2025). The study contributes to the design of insurance mechanisms that encourage policyholders' prevention efforts and maximize the welfare of both insurers and policyholders, thus promoting proactive risk management and enhanced market resilience.

Economically, unless a loss reporting stage is explicitly modeled, there may be little room for true 'incentives'. In our model, information about effort is observable and contractible, and manipulation or false reporting are not considered. As a result, incentive compatibility and the no-sabotage condition are trivially satisfied. However, in a more general setting where only the insured observes the actual loss and may choose to misreport its value, the incentive effect of cost-sharing could weaken significantly. We argue that risk sharing offers a potential solution to this issue. Hu et al. (2025) examine the role of risk sharing in mitigating ex ante moral hazard in P2P insurance. The study notes that while risk sharing alleviates moral hazard, it may also lead to free-riding in effort provision.

Several promising research directions remain for future exploration.

First, the assumption of proportional insurance could be relaxed in favor of a more general indemnity structure that satisfies the incentive-compatibility condition proposed by Huberman et al. (1983). In such a framework, proportional insurance may no longer retain its optimality, opening avenues to analyze alternative contractual forms.

Second, while our model assumes full observability of the insured's prevention effort by the insurer, this premise may conflict with real-world scenarios where effort levels are privately held information. Extending the analysis to incorporate ex ante moral hazard, where the insured selects prevention efforts contingent on the insurance contract, could yield valuable insights into the interplay between insurance demand, self-protection, and cost-sharing mechanisms (Parra & Winter, 2013; Seog, 2012; Seog & Hong, 2024).

A third direction involves examining a Stackelberg game involving one insurer and multiple insureds whose losses or protective measures exhibit interdependencies. Such a setting could shed light on strategic interactions under risk correlation, building on frameworks proposed by Muermann and Kunreuther (2008) and Hofmann and Peter (2015).

Acknowledgments

The authors are grateful to two anonymous referees for their constructive comments and suggestions, which have significantly improved the paper.

Disclosure statement

No potential conflict of interest was reported by the author(s).

Funding

Yiying Zhang acknowledges the financial support from the National Natural Science Foundation of China (No. 12571513), and Shenzhen Science and Technology Program (No. JCYJ20230807093312026 and No. JCYJ20250604144210013).

References

- Altintas, N., Erhun, F., & Tayur, S. (2008). Quantity discounts under demand uncertainty. *Management Science*, 54 (4), 777–792. <https://doi.org/10.1287/mnsc.1070.0829>
- Augeraud-Véron, E., & Leandri, M. (2024). Optimal self-protection and health risk perceptions: Exploring connections between risk theory and the health belief model. *Health Economics*, 33(7), 1565–1583. <https://doi.org/10.1002/hec.v33.7>
- Bensalem, S., Santibáñez, N. H., & Kazi-Tani, N. (2020). Prevention efforts, insurance demand and price incentives under coherent risk measures. *Insurance: Mathematics and Economics*, 93, 369–386.
- Cao, J., Li, D., Young, V. R., & Zou, B. (2023). Reinsurance games with variance-premium reinsurers: From tree to chain. *ASTIN Bulletin: The Journal of the IAA*, 53(3), 706–728. <https://doi.org/10.1017/asb.2023.24>

- Cao, J., Li, D., Young, V. R., & Zou, B. (2024a). Short communication: Optimal insurance to maximize exponential utility when premium is computed by a convex functional. *SIAM Journal on Financial Mathematics*, 15(1), SC15–SC27. <https://doi.org/10.1137/23M1601237>
- Cao, J., Li, D., Young, V. R., & Zou, B. (2024b). Stackelberg reinsurance chain under model ambiguity. *Scandinavian Actuarial Journal*, 2024(4), 329–360. <https://doi.org/10.1080/03461238.2023.2255399>
- Chase, J., Niyato, D., Wang, P., Chaisiri, S., & Ko, R. K. (2017). A scalable approach to joint cyber insurance and security-as-a-service provisioning in cloud computing. *IEEE Transactions on Dependable and Secure Computing*, 16(4), 565–579. <https://doi.org/10.1109/TDSC.8858>
- Chen, L., Zhang, Y., & Zhu, X. (2026). Self-protection and insurance demand under time-inconsistent mean–variance framework. *European Journal of Operational Research*, 330(1), 199–216. <https://doi.org/10.1016/j.ejor.2025.08.004>
- Chen, S., Zhang, Y., & Zhu, X. (2025). Interplay of prevention efforts and insurance demand with distortion risk measures. Available at SSRN 5145103.
- Chen, Z., Chen, B., & Mao, Y. (2024). Fence off black swans: The economics of insurance for vaccine injury. *Oxford Bulletin of Economics and Statistics*, 86(5), 995–1025. <https://doi.org/10.1111/obes.v86.5>
- Chen, Z., Dhaene, J., & Li, T. (2024). Moral hazard in peer-to-peer insurance with social connection. Available at SSRN 4846469.
- Chi, Y., Peter, R., & Qi, M. (2025). Disentangling risk and time for optimal prevention. Available at SSRN 5248267.
- Chiu, W. H. (2005). Degree of downside risk aversion and self-protection. *Insurance: Mathematics and Economics*, 36(1), 93–101.
- Corbett, C. J., & De Groote, X. (2000). A supplier's optimal quantity discount policy under asymmetric information. *Management Science*, 46(3), 444–450. <https://doi.org/10.1287/mnsc.46.3.444.12065>
- Courbage, C. (2001). Self-insurance, self-protection and market insurance within the dual theory of choice. *The Geneva Papers on Risk and Insurance Theory*, 26(1), 43–56. <https://doi.org/10.1023/A:1011212324117>
- Courbage, C., & Coulon, A. d. (2004). Prevention and private health insurance in the UK. *The Geneva Papers on Risk and Insurance-Issues and Practice*, 29(4), 719–727. <https://doi.org/10.1111/j.1468-0440.2004.00313.x>
- Courbage, C., Loubergé, H., & Peter, R. (2017). Optimal prevention for multiple risks. *Journal of Risk and Insurance*, 84(3), 899–922. <https://doi.org/10.1111/jori.v84.3>
- Crainich, D., & Eeckhoudt, L. (2017). Average willingness to pay for disease prevention with personalized health information. *Journal of Risk and Uncertainty*, 55(1), 29–39. <https://doi.org/10.1007/s11666-017-9265-z>
- Crainich, D., & Menegatti, M. (2021). Self-protection with random costs. *Insurance: Mathematics and Economics*, 98, 63–67.
- Cui, G.-H., Wang, Z., Li, J.-L., Jin, X., & Zhang, Z.-W. (2021). Influence of precaution and dynamic post-indemnity based insurance policy on controlling the propagation of epidemic security risks in networks. *Applied Mathematics and Computation*, 392, 125720. <https://doi.org/10.1016/j.amc.2020.125720>
- Denuit, M. M., Eeckhoudt, L., Liu, L., & Meyer, J. (2016). Tradeoffs for downside risk-averse decision-makers and the self-protection decision. *The Geneva Risk and Insurance Review*, 41(1), 19–47. <https://doi.org/10.1057/grir.2015.3>
- Deprez, O., & Gerber, H. U. (1985). On convex principles of premium calculation. *Insurance: Mathematics and Economics*, 4(3), 179–189.
- Ehrlich, I., & Becker, G. S. (1972). Market insurance, self-insurance, and self-protection. *Journal of Political Economy*, 80(4), 623–648. <https://doi.org/10.1086/259916>
- He, L., Hou, P., & Liang, Z. (2008). Optimal control of the insurance company with proportional reinsurance policy under solvency constraints. *Insurance: Mathematics and Economics*, 43(3), 474–479.
- Hofmann, A., & Peter, R. (2015). Multivariate prevention decisions: Safe today or sorry tomorrow? *Economics Letters*, 128, 51–53. <https://doi.org/10.1016/j.econlet.2015.01.016>
- Hofmann, A., & Rothschild, C. (2019). On the efficiency of self-protection with spillovers in risk. *The Geneva Risk and Insurance Review*, 44(2), 207–221. <https://doi.org/10.1057/s10713-019-00041-z>
- Hu, W., Chen, Z., Ma, W., & Shi, Y. (2024). Optimal usage strategy of catastrophe insurance compensation based on epidemiological model. Available at SSRN 4974197.
- Hu, W., Li, C., Ling, B., & Zhang, Y. (2025). Optimal self-protection in peer-to-peer insurance: The dual role of risk sharing. Available at SSRN 5583170.
- Huang, C.-Y., Yang, M.-C., Huang, C.-Y., Chiu, P.-S., Liu, Z.-S., & Chang, R.-I. (2017). Design and implementation of a dynamic healthcare system for weight management and health promotion. In *2017 IEEE International Conference on Industrial Engineering and Engineering Management (IEEM)* (pp. 2386–2390). IEEE.
- Huberman, G., Mayers, D., & Smith Jr, C. W. (1983). Optimal insurance policy indemnity schedules. *Bell Journal of Economics*, 14(2), 415–426. <https://doi.org/10.2307/3003643>
- Kaluszka, M. (2005). Optimal reinsurance under convex principles of premium calculation. *Insurance: Mathematics and Economics*, 36(3), 375–398.
- Kelly, M., & Kleffner, A. E. (2003). Optimal loss mitigation and contract design. *Journal of Risk and Insurance*, 70(1), 53–72. <https://doi.org/10.1111/jori.2003.70.issue-1>
- Khan, S., & Yairi, T. (2018). A review on the application of deep learning in system health management. *Mechanical Systems and Signal Processing*, 107, 241–265. <https://doi.org/10.1016/j.ymssp.2017.11.024>

- Klimaviciute, J. (2017). Long-term care insurance and intra-family moral hazard: Fixed vs proportional insurance benefits. *The Geneva Risk and Insurance Review*, 42(2), 87–116. <https://doi.org/10.1057/s10713-016-0018-8>
- Lee, H. L., & Rosenblatt, M. J. (1986). A generalized quantity discount pricing model to increase supplier's profits. *Management Science*, 32(9), 1177–1185. <https://doi.org/10.1287/mnsc.32.9.1177>
- Li, J., Peter, R., & Zhou, L. (2025). Cross-prudence and optimal prevention. Available at SSRN 5274609.
- Li, L., & Peter, R. (2021). Should we do more when we know less? The effect of technology risk on optimal effort. *Journal of Risk and Insurance*, 88(3), 695–725. <https://doi.org/10.1111/jori.v88.3>
- Li, Q., Wang, W., & Zhang, Y. (2025). Self-protection and insurance demand with convex premium principles. *North American Actuarial Journal*, 1–35. <https://doi.org/10.1080/10920277.2025.2561055>
- Lin, M., Xu, L., Yang, L. T., Qin, X., Zheng, N., Wu, Z., & Qiu, M. (2009). Static security optimization for real-time systems. *IEEE Transactions on Industrial Informatics*, 5(1), 22–37. <https://doi.org/10.1109/TII.2009.2014055>
- Macé, S., & Peter, R. (2021). The impact of loss aversion on optimal prevention. Available at SSRN 3882200.
- Madubuike, O. C., Anumba, C. J., & Agapaki, E. (2024). Scenarios for digital twin deployment in healthcare facilities management. *Journal of Facilities Management*, 22(5), 900–919. <https://doi.org/10.1108/JFM-10-2022-0107>
- Menegatti, M. (2014). Optimal choice on prevention and cure: A new economic analysis. *The European Journal of Health Economics*, 15(4), 363–372. <https://doi.org/10.1007/s10198-013-0479-y>
- Meyer, D. J., & Meyer, J. (2011). A diamond-stiglitz approach to the demand for self-protection. *Journal of Risk and Uncertainty*, 42(1), 45–60. <https://doi.org/10.1007/s11166-010-9107-8>
- Monahan, J. P. (1984). A quantity discount pricing model to increase vendor profits. *Management Science*, 30(6), 720–726. <https://doi.org/10.1287/mnsc.30.6.720>
- Muermann, A., & Kunreuther, H. (2008). Self-protection and insurance with interdependencies. *Journal of Risk and Uncertainty*, 36(2), 103–123. <https://doi.org/10.1007/s11166-008-9033-1>
- Peter, R. (2017). Optimal self-protection in two periods: On the role of endogenous saving. *Journal of Economic Behavior and Organization*, 137, 19–36. <https://doi.org/10.1016/j.jebo.2017.01.017>
- Peter, R. (2021a). A fresh look at primary prevention for health risks. *Health Economics*, 30(5), 1247–1254. <https://doi.org/10.1002/hec.v30.5>
- Peter, R. (2021b). Prevention as a giffen good. *Economics Letters*, 208, 110052. <https://doi.org/10.1016/j.econlet.2021.110052>
- Peter, R. (2021c). Who should exert more effort? Risk aversion, downside risk aversion and optimal prevention. *Economic Theory*, 71(4), 1259–1281. <https://doi.org/10.1007/s00199-020-01282-0>
- Peter, R. (2024b). The economics of self-protection: A tribute to EGRIE. *The Geneva Risk and Insurance Review*, 49(1), 6–35. <https://doi.org/10.1057/s10713-023-00094-1>
- Peter, R. (2025). The Self-Protection Problem. In G. Dionne (Ed.), *Handbook of insurance: Volume II* (pp. 55–82). Cham: Springer. https://doi.org/10.1007/978-3-031-69674-9_3
- Ram, C. P., & Sreenivaasan, G. (2010). Security as a service (sass): securing user data by coprocessor and distributing the data. In *Trendz in Information Sciences and Computing (TISC2010)* (pp. 152–155). IEEE.
- Seog, S. H. (2012). Moral hazard and health insurance when treatment is preventive. *Journal of Risk and Insurance*, 79(4), 1017–1038. <https://doi.org/10.1111/jori.2012.79.issue-4>
- Seog, S. H., & Hong, J. (2024). Moral hazard in loss reduction and state-dependent utility. *Insurance: Mathematics and Economics*, 115, 151–168.
- Shaanxi (2025). Announcement from the shaanxi provincial emergency management department. Shaanxi Provincial Emergency Management Department official website. https://yjtt.shaanxi.gov.cn/gk/zcwj/wjzl/gggs/202508/t20250815_3555472.html
- Wang, Q., Egelandsdal, B., Amdam, G. V., Almli, V. L., & Oostindjer, M. (2016). Diet and physical activity apps: Perceived effectiveness by app users. *JMIR MHealth and UHealth*, 4(2), e33. <https://doi.org/10.2196/mhealth.5114>
- Weng, Z. K. (1995). Channel coordination and quantity discounts. *Management Science*, 41(9), 1509–1522. <https://doi.org/10.1287/mnsc.41.9.1509>
- Williams, J., Degeling, C., McVernon, J., & Dawson, A. (2021). How should we conduct pandemic vaccination? *Vaccine*, 39(6), 994–999. <https://doi.org/10.1016/j.vaccine.2020.12.059>
- Winter, R.A. (2013). Optimal insurance contracts under moral hazard. In G. Dionne (Ed.), *Handbook of insurance: Volume II* (pp. 129–164). New York, NY: Springer. https://doi.org/10.1007/978-1-4614-0155-1_9
- Wu, C., & Colwell, P. F. (1988). Moral hazard and moral imperative. *Journal of Risk and Insurance*, 55(1), 101–117. <https://doi.org/10.2307/253283>
- Yin, Y., & Meng, S. (2025). Self-protection under n th-degree risk increase of random unit cost. *Insurance: Mathematics and Economics*, 122, 137–142.
- Zeller, G., & Scherer, M. (2023). Risk mitigation services in cyber insurance: Optimal contract design and price structure. *The Geneva Papers on Risk and Insurance – Issues and Practice*, 48(2), 502–547. <https://doi.org/10.1057/s41288-023-00289-7>
- Zweifel, P., Eisen, R., Zweifel, P., & Eisen, R. (2012). Insurance demand I: Decisions under risk without diversification possibilities. In *Insurance economics* (pp. 71–110). Springer.

Appendix

A.1 Optimal indemnity function under mean-variance preference

Theorem A.1: Under the mean-variance objective given in Problem 2.1, the optimal indemnity function $I^*(x)$ is given by

$$I^*(x) = \frac{(\gamma x + \lambda^*)_+}{\theta_2 + \gamma}, \quad (\text{A1})$$

where $\lambda^* < 0$ uniquely solves

$$\lambda = \gamma \left(\frac{\gamma}{\gamma + \theta_2} \int_{-\lambda/\gamma}^{+\infty} S_X(x) dx - \mu \right) + 1 - \theta_1. \quad (\text{A2})$$

Proof: Let the indemnity function be $I(x)$, then the premium is

$$\pi(I(X)) = \frac{\theta_2}{2} \mathbb{E}[I^2(X)] + \theta_1 \mathbb{E}[I(X)].$$

Then, the wealth of insured is

$$W_b = w - \pi(I(X)) - X + I(X) - c(e).$$

The optimization problem is

$$I^*(x) = \arg \max_{0 \leq I(x) \leq x} \mathbb{E}[W_b] - \frac{\gamma}{2} \text{Var}[W_b]. \quad (\text{A3})$$

By expanding $\mathbb{E}[W_b] - \frac{\gamma}{2} \text{Var}[W_b]$, one can see that solving (A3) is equivalent to maximizing $J(I)$, in which

$$J(I) = -\frac{\theta_2 + \gamma}{2} \mathbb{E}[I^2] + \gamma \mathbb{E}[XI] + \frac{\gamma}{2} \mathbb{E}^2[I] - (\theta_1 - 1) \mathbb{E}[I] - \gamma \mu \mathbb{E}[I].$$

To maximize $J(I)$, first fix $\mathbb{E}[I] = m \in [0, \mu]$, where $\mu = \mathbb{E}[X]$. Then, maximizing $J(I)$, subject to $\mathbb{E}[I] = m$ and $0 \leq I(X) \leq X$ is equivalent to maximizing

$$-\frac{\theta_2 + \gamma}{2} \mathbb{E}[I^2] + \gamma \mathbb{E}[XI]$$

subject to $\mathbb{E}[I] = m$ and $0 \leq I(X) \leq X$. To that end, we define the Lagrangian L by

$$L(I) = -\frac{\theta_2 + \gamma}{2} \mathbb{E}[I^2] + \gamma \mathbb{E}[XI] + \lambda(\mathbb{E}[I] - m),$$

where λ is the Lagrange multiplier. We rewrite $L(I)$ as follows

$$L(I) = \int_0^{+\infty} \left[-\frac{\theta_2 + \gamma}{2} I^2 + \gamma xI + \lambda I \right] dF_X(x) - \lambda m.$$

We can maximize $L(I)$ by maximizing the integrand x -by- x , subject to $0 \leq I(x) \leq x$. Thus, we can obtain

$$I^*(x) = \min \left\{ \frac{(\gamma x + \lambda)_+}{\theta_2 + \gamma}, x \right\}.$$

Next, we show that, given $m \in [0, \mu]$, there exists a unique value of λ such that

$$m = \mathbb{E} \left[\min \left\{ \frac{(\gamma x + \lambda)_+}{\theta_2 + \gamma}, x \right\} \right]. \quad (\text{A4})$$

It is easy to obtain that when $\lambda \geq 0$, we have

$$m = \int_0^{\frac{\lambda}{\theta_2}} x dF_X(x) + \int_{\frac{\lambda}{\theta_2}}^{+\infty} \frac{\gamma x + \lambda}{\theta_2 + \gamma} dF_X(x),$$

and when $\lambda < 0$, we have

$$m = \int_{-\frac{\lambda}{\gamma}}^{+\infty} \frac{\gamma x + \lambda}{\theta_2 + \gamma} dF_X(x).$$

It is easy to check that m is increasing with respect to λ , and increases from 0 to μ when λ increases from $-\infty$ to $+\infty$. It follows that (A4) has a unique solution λ . Thus, we can maximize $J(I)$ by finding the maximizing value of λ .

We firstly consider the case where $\lambda > 0$. In this case, we can obtain:

$$J'(\lambda) = (\gamma m - \theta_1 + 1 - \gamma \mu - \lambda) \int_{\frac{\lambda}{\theta_2}}^{+\infty} \frac{1}{\theta_2 + \gamma} f(x) dx.$$

It is easy to check that $J'(\lambda) < 0$ in this case, which implies that the optimal λ is non-positive. For $\lambda \leq 0$, we obtain

$$J'(\lambda) = (\gamma m - \theta_1 + 1 - \gamma \mu - \lambda) \int_{-\frac{\lambda}{\gamma}}^{+\infty} \frac{1}{\theta_2 + \gamma} f(x) dx.$$

Define $g(\lambda) = \gamma m - \theta_1 + 1 - \gamma \mu - \lambda$, we obtain

$$g'(\lambda) = \frac{\gamma}{\theta_2 + \gamma} S_X\left(-\frac{\lambda}{\gamma}\right) - 1 < 0.$$

Thus, it is easy to see that there uniquely exists a $\lambda^* < 0$ that solves $g(\lambda) = 0$ such that $J'(\lambda)$ is positive for $\lambda < \lambda^*$ and negative for $\lambda > \lambda^*$. $g(\lambda) = 0$ implies

$$\gamma \int_{-\frac{\lambda}{\gamma}}^{+\infty} \frac{\gamma x + \lambda}{\theta_2 + \gamma} f(x) dx - \theta_1 + 1 - \gamma \mu - \lambda = 0,$$

which can be rewritten as

$$\lambda = \gamma \left(\frac{\gamma}{\gamma + \theta_2} \int_{-\frac{\lambda}{\gamma}}^{+\infty} S_X(x) dx - \mu \right) + 1 - \theta_1.$$

Then, we have proved this theorem. ■

As shown in Theorem A.1, optimal coverage reduces to $I^* = x/\theta_2$ when $\theta_1 = 1$ and $\gamma = 0$. We thoroughly investigated the possibility of deriving optimal parameters for the general form defined by (A1). However, the resulting mathematical complexity precludes a semi-explicit solution. To ensure analytical tractability and to facilitate a clear presentation of the key economic insights, we have therefore chosen to work with the proportional insurance assumption in the main analysis. We believe this simplification strikes a balance between generality and clarity.

A.2 Results under the simplified cost-sharing mechanism

We note that the cost-sharing mechanisms in Cui et al. (2021) and Shaanxi (2025) are special cases of our proposed framework. Specifically, Cui et al. (2021) assumes that when the insured purchases insurance, the insurer covers a fixed proportion of the effort cost. In the context of our model, this corresponds to defining $\beta \in [0, 1]$ if $a > 0$, and $\beta = 1$ if $a = 0$. Shaanxi (2025) omits price discrimination, which corresponds to $\beta_0 = 1$ in our setting. In what follows, we analyze a simpler and more transparent model with $\beta \in [0, 1]$ if $a > 0$, $\beta = 1$ if $a = 0$, and $\beta_0 = 1$.

We first study Problem (11), to determine the optimal insurance demand $a^*(e)$ given a fixed e . The result is as follows:

Proposition A.2: For any given effort $e \in \mathbb{R}_+$, the optimal insurance demand $a^*(e)$ is:

- If $e = 0$, then $a^*(e) = 0$.
- If $e \in (\bar{e}, +\infty)$, where $\bar{e}$ uniquely solves:

$$\frac{\mathbb{E}[X_e]}{\text{Var}[X_e]} = \frac{\gamma_1}{\theta_1 - 1},$$

then $a^*(e)$ is given by:

$$a^*(e) = \frac{-(\theta_1 - 1)\mathbb{E}[X_e] + \gamma_1 \text{Var}[X_e]}{\theta_2 a \mathbb{E}[X_e]^2 + \gamma_1 \text{Var}[X_e]}. \quad (\text{A5})$$

- If $e \in (0, \bar{e}]$, $a^*(e)$ does not exist.

Proof: Under the simplified cost-sharing mechanism, the objective function is defined as

$$L_e(a) = (1 - a)\mathbb{E}[X_e] + \frac{\theta_2 a^2}{2} \mathbb{E}[X_e]^2 + \theta_1 a \mathbb{E}[X_e] + \mathbf{1}_{(0,1]}(a) \beta c(e) + \mathbf{1}_{\{0\}}(a) c(e) + \frac{\gamma_1}{2} \text{Var}[(1 - a)X_e].$$

For $a = 0$, the expression reduces to

$$L_e(0) = \mathbb{E}[X_e] + \frac{\gamma_1}{2} \text{Var}[X_e] + c(e).$$

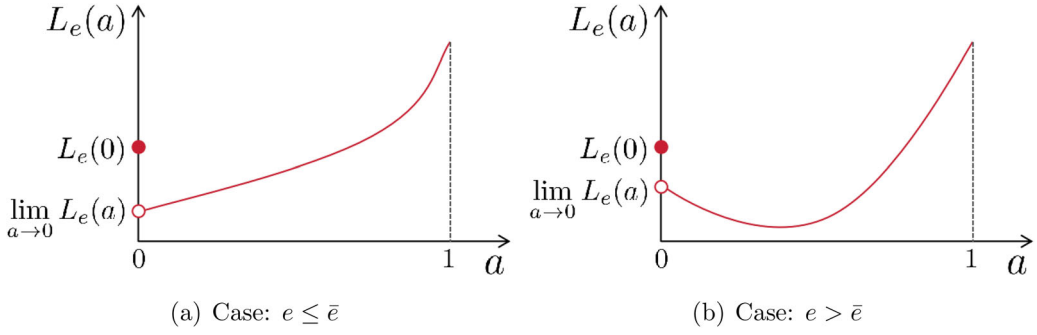


Figure A1. Determination of $a^*(e)$ under simple cost-sharing mechanism. (a) Case: $e \leq \bar{e}$ and (b) Case: $e > \bar{e}$.

For $a > 0$, it becomes

$$L_e(a) = (1-a)\mathbb{E}[X_e] + \frac{\theta_2 a^2}{2} \mathbb{E}[X_e]^2 + \theta_1 a \mathbb{E}[X_e] + \beta c(e) + \frac{\gamma_1}{2} \text{Var}[(1-a)X_e].$$

The first and second derivatives of $L_e(a)$ are

$$\begin{aligned} L'_e(a) &= (\theta_1 - 1)\mathbb{E}[X_e] + \theta_2 a \mathbb{E}[X_e]^2 - (1-a)\gamma_1 \text{Var}[X_e]. \\ L''_e(a) &= \theta_2 \mathbb{E}[X_e]^2 + \gamma_1 \text{Var}[X_e] > 0. \end{aligned}$$

Since $L''_e(a) > 0$, $L'_e(a)$ is strictly increasing in a . Then, we can obtain

$$\begin{aligned} \max L'_e(a) &= L'_e(1) = (\theta_1 - 1)\mathbb{E}[X_e] + \theta_2 \mathbb{E}[X_e]^2 > 0. \\ \min L'_e(a) &= L'_e(0) = (\theta_1 - 1)\mathbb{E}[X_e] - \gamma_1 \text{Var}[X_e]. \end{aligned}$$

Therefore,

- If $L'_e(0) \geq 0$, then $L_e(a)$ is increasing in a ;
- If $L'_e(0) < 0$, then $L_e(a)$ first decreases and then increases in a .

Given that $\frac{\mathbb{E}[X_e]}{\text{Var}[X_e]}$ is decreasing in e , the condition $L'_e(0) \geq 0$ is equivalent to $e \leq \bar{e}$, where $\bar{e}$ is the unique solution to

$$\frac{\mathbb{E}[X_e]}{\text{Var}[X_e]} = \frac{\gamma_1}{\theta_1 - 1}.$$

To avoid trivial cases, we assume such $\bar{e}$ always exists. Correspondingly, $L'_e(0) < 0$ is equivalent to $e > \bar{e}$.

Now consider the case where $e \leq \bar{e}$. Since $L_e(a)$ is discontinuous at $a = 0$, and

$$\lim_{a \rightarrow 0} L_e(a) = \mathbb{E}[X_e] + \frac{\gamma_1}{2} \text{Var}[X_e] + \beta c(e) \leq L_e(0),$$

a solution $a^*(e) = 0$ exists only when $e = 0$, in which case the limit $\lim_{a \rightarrow 0} L_e(a)$ equals $L_e(0)$. For $0 < e \leq \bar{e}$, no solution exists. These situation is illustrated in panel (a) of Figure A1.

Next, consider the case where $e > \bar{e}$. Here, the optimal demand $a^*(e)$ lies in $(0, 1)$ and satisfies the first-order condition $L'_e(a) = 0$, which yields:

$$a^*(e) = \frac{-(\theta_1 - 1)\mathbb{E}[X_e] + \gamma_1 \text{Var}[X_e]}{\theta_2 \mathbb{E}[X_e]^2 + \gamma_1 \text{Var}[X_e]}.$$

The determination of $a^*(e)$ in this case is shown in panel (b) of Figure A1. ■

We now characterize the monotonicity of the optimal demand $a^*(e)$ with respect to e .

Proposition A.3: $a^*(e)$ is increasing with e .

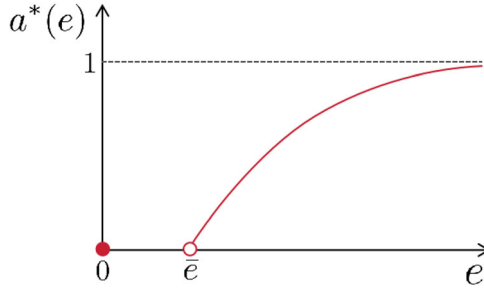


Figure A2. $a^*(e)$ under simple cost-sharing mechanism.

Proof: We now analyze the monotonicity of $a^*(e)$ with e when it is determined by Equation (A5). We begin by rewriting the expression as:

$$a^*(e) = \frac{-\frac{\theta_1 - 1}{\gamma_1} \frac{\mathbb{E}[X_e]}{\text{Var}[X_e]} + 1}{\frac{\theta_2}{\gamma_1} \frac{\mathbb{E}[X_e]^2}{\text{Var}[X_e]} + 1}.$$

It is easy to check that $\frac{\mathbb{E}[X_e]^2}{\text{Var}[X_e]}$ is decreasing in e , thus, $a^*(e)$ is increasing in e . ■

Proposition A.2 characterizes the optimal demand $a^*(e)$ as follows: it is zero when $e = 0$, does not exist for $e \in (0, \bar{e}]$, and uniquely exists in $(0, 1)$ for $e > \bar{e}$, as given by Equation (A5). Proposition A.3 establishes that $a^*(e)$ is increasing in e . Furthermore, as $e \rightarrow \bar{e}$, we observe that $a^*(e) \rightarrow 0$.

Two key distinctions emerge compared to $a^*(e)$ we obtained in Section 3.1. First, in the simplified model, $a^*(e)$ does not always exist over the entire domain of e . This arises from the discontinuity in the insurer's share of the effort cost at $a = 0$: when the insured purchases insurance, the cost-sharing ratio jumps from 0 to β . Second, $a^*(e)$ is always strictly less than 1. This occurs because the insurer's shared effort cost does not vary with the number of policies a when $a > 0$. Consequently, the insured has no incentive to reduce their effort cost by buying more insurance. The shape of $a^*(e)$ is illustrated in Figure A2.

A.3 Some examples satisfying Assumption thm3.1

Now, let us examine the conditions under which Assumption 3.1 holds. By observing the expression for $H(e)$ in Equation (14), we find that the expectation term always satisfies the following conditions:

$$\frac{d\mathbb{E}[X_e]}{de} < 0 \quad \text{and} \quad \frac{d^2\mathbb{E}[X_e]}{de^2} > 0.$$

Thus, one sufficient condition to ensure Assumption 3.1 to hold depends primarily on the behavior of the variance term, $\text{Var}[X_e]$. If γ_1 is small, Assumption 3.1 is likely satisfied. However, to allow γ_1 to take any value in $\mathbb{R}_+$, we examine sufficient conditions for Assumption 3.1, which are:

$$\frac{d\text{Var}[X_e]}{de} \leq 0 \quad \text{and} \quad \frac{d^2\text{Var}[X_e]}{de^2} \geq 0.$$

These conditions are equivalent to:

$$\begin{aligned} p'(e)[\sigma_0^2 + (1 - 2p(e))\mu_0^2] &\leq 0, \\ p''(e)[\sigma_0^2 + (1 - 2p(e))\mu_0^2] - 2p'(e)^2\mu_0^2 &\geq 0. \end{aligned}$$

Solving these inequalities leads to the following sufficient condition for Assumption 3.1:

$$\frac{\sigma_0^2}{\mu_0^2} \geq \frac{2p'(e)^2}{p''(e)} + 2p(e) - 1.$$

Here, we consider an exponentially decaying form of $p(e)$, analyzed under both the Gamma and Pareto distributions. Let

$$p(x) = (\bar{p} - \underline{p})e^{-x} + \underline{p}. \tag{A6}$$

Then we have:

$$\frac{2p'(e)^2}{p''(e)} + 2p(e) - 1 \leq 4\bar{p} - 2\underline{p} - 1.$$

Thus, Assumption 3.1 holds if:

$$\frac{\sigma_0^2}{\mu_0^2} \geq 4\bar{p} - 2\underline{p} - 1.$$

Consider the Gamma distribution $\Gamma(k, \theta)$ where $\mu_0 = k\theta$ and $\sigma_0^2 = k\theta^2$, we can obtain that Assumption 3.1 holds if:

$$\frac{1}{k} \geq 4\bar{p} - 2\underline{p} - 1.$$

For the Pareto distribution $P(k, x_m)$, we have $\mu_0 = \frac{kx_m}{k-1}$ and $\sigma_0^2 = \frac{kx_m^2}{(k-1)^2(k-2)}$. Hence, Assumption 3.1 holds if:

$$\frac{1}{k(k-2)} \geq 4\bar{p} - 2\underline{p} - 1.$$

In the following analysis, we consider a rational form of $p(e)$, examining its behavior under both the Gamma and Pareto distributions. Let

$$p(x) = \frac{1}{\frac{1}{\bar{p}-\underline{p}} + x} + \underline{p}.$$

In this case, Assumption 3.1 holds if:

$$\frac{\sigma_0^2}{\mu_0^2} \geq 3\bar{p} - \underline{p} - 1.$$

Then, for the Gamma distribution $\Gamma(k, \theta)$, Assumption 3.1 holds if:

$$\frac{1}{k} \geq 3\bar{p} - \underline{p} - 1.$$

For the Pareto distribution $P(k, x_m)$, Assumption 3.1 holds if:

$$\frac{1}{k(k-2)} \geq 3\bar{p} - \underline{p} - 1.$$