

Adverse Selection, Belief Heterogeneity and Market Power in Insurance Markets

Karl Whelan*

University College Dublin and CEPR

September 2026

Abstract

Increased risk heterogeneity is associated with adverse selection problems that lead to higher insurance premiums and the possibility that markets collapse. This paper shows how these properties rely on consumers knowing their probability of loss and insurers being price takers. We examine how adverse selection, heterogeneous beliefs about loss probability, and market power interact in a pooled insurance market. Consumers' beliefs about their loss probabilities equal their true values plus an idiosyncratic component. We present a series of model specifications for this component, and for dispersion in objective risk, that generate unconventional results. Greater adverse selection can leave premiums unchanged or even reduce them. We also describe examples in which, at higher levels of risk heterogeneity, rather than collapse, markets continue to exist but gradually stop serving lower-risk consumers. This framework is used to show how policies to replace risk-based pricing with pooled pricing can increase social welfare.

Keywords: Insurance, adverse selection, pooled contracts, belief heterogeneity, market power

JEL Classification: D42, D82, G22

*Email: karl.whelan@ucd.ie.

1. Introduction

Since Akerlof (1970), adverse selection problems in insurance markets have been associated with higher premiums and the possibility that markets collapse. This paper discusses how these results can change when we weaken the assumptions that consumers know their probability of loss and that insurers are price takers. Neither assumption holds in real-world insurance markets. The evidence, as presented by Sydnor (2010), Barseghyan et al. (2013) and Abito and Salant (2019), shows that consumers hold heterogeneous subjective beliefs about their probability of loss, distinct from their objective risk. Indeed, in an application to auto insurance markets, Cohen and Einav (2007) show idiosyncratic components account for more of the variation in insurance demand than heterogeneity in objective risk. Insurance markets also typically feature suppliers with market power, as discussed by Stiglitz (1977) and Mahoney and Weyl (2017). In some cases, this arises from natural monopolies such as manufacturers selling warranties on their own products but, more generally, the highly regulated nature of insurance markets tends to reduce competition.

We use a pooled insurance framework to show that once customers with identical objective loss probabilities have heterogeneous beliefs about this risk, and insurers have market power, the standard results about adverse selection can be reversed. We show how, within realistic ranges for risk and belief heterogeneity, greater adverse selection can leave premiums unchanged or even reduce them. We also describe examples in which, for higher levels of risk heterogeneity, rather than collapse, markets continue existing but gradually stop serving lower-risk consumers.

The pooled full-insurance environment studied here is a natural one. Contracts like this are common: extended warranties, smartphone protection plans, travel insurance policies, and many health and car insurance products are sold through relatively simple pooled pricing structures despite substantial heterogeneity in risks and beliefs. In addition, regulations such as the EU's ban on gender-based pricing and restrictions on the use of genetic or medical history information in health and life insurance, often mandate pooled pricing. More generally, while the insurance literature has typically focused on insurers responding to adverse selection by offering a wide menu of screening contracts, as in Rothschild and Stiglitz (1976), practicality often requires offering a limited number of contracts, each of which will be purchased by consumers with varying levels of risk.

Model features and results

Our framework features a continuum of customers with varying objective risk of loss. Each customer has a subjective belief about their probability of loss that equals its true value plus an idiosyncratic component. We present two approaches to modeling the idiosyncratic element of beliefs and consider three different approaches to modeling risk heterogeneity.

Our initial implementation makes the form of risk heterogeneity very general, assuming only that belief heterogeneity is large enough so that the fraction of all risk groups that buys the insurance lies

between zero and one at the profit-maximizing premium. In this case, when individual beliefs deviate from the group's objective probability by a random amount drawn from the same distribution for every risk group, demand curves for different risk groups are parallel shifts of one another. This implies aggregate demand is identical to that of a homogeneous population with the same mean risk, which means the extra claims generated by a riskier pool can act like a fixed reduction in profit, but do not change profit-maximizing premium. However, sufficiently high risk heterogeneity can still lead to the market collapsing: the market either exists at a fixed price or else it collapses.

Since the extent of belief heterogeneity about risk likely scales with actual risk probabilities, the paper focuses more on a second specification, in which consumers who face higher objective risk also disagree more, in absolute terms, about their own probability of loss. Sydnor (2010) and Abito and Salant (2019) both find that overestimation of claim rates scales with objective risk rather than remaining constant across risk groups. In this case, demand curves for different risk groups are not parallel: higher-risk consumers hold wider absolute distributions of belief and this makes them less price-sensitive than lower-risk consumers. This means aggregate demand is steeper than in the homogeneous benchmark in which all consumers have the same loss probability. As a result, adverse selection can lower premiums relative to this benchmark. Again, however, sufficiently high risk heterogeneity can lead to the monopoly insurance market collapsing. We explain the importance of market structure for these findings by showing how the traditional result of adverse selection worsening prices still holds for this model with perfect competition and that premiums in a Cournot equilibrium tend toward competitive levels as the number of insurers increases.

We then present a specific model of risk heterogeneity — objective risk varies across customers according to a factor of proportionality drawn from a uniform distribution — that allows for an analytical solution. If risk heterogeneity is large enough, some low-risk groups decide not to purchase any insurance but we find that the market does not collapse. Once lower-risk groups start to leave the market, higher risk heterogeneity begins to raise optimal premiums and the market survives by charging these higher premiums to an ever-worsening risk pool. We show how this result also applies to Cournot oligopoly when the number of insurers is small, though the rising premium result typically occurs at higher levels of risk heterogeneity than are likely seen in real insurance markets.

Our last model of risk heterogeneity assumes a truncated Normal distribution for risk probabilities and beliefs about them, truncated from below so negative risk probabilities or beliefs are not possible. This does not have an analytical solution but numerical calculations point to similar unconventional conclusions. Premiums still fall slightly as risk heterogeneity rises from low levels — though by less than under uniformly distributed risk and beliefs — before rising but, again, without the market collapsing.

Our base case assumes the insurer offers a full insurance contract. We also consider the case where a monopolist can choose coinsurance jointly with the premium, and find that full coverage

departs from optimality only when loss size and risk aversion are simultaneously pushed well beyond plausible values.

Finally, the paper also derives the consequences of replacing risk-based pricing — in which the insurer can observe each risk group’s identity and charge group-specific premiums — with pooled pricing. With proportionally scaled beliefs, pooled pricing can increase total welfare under monopoly, because the insurer’s ability to extract surplus from high-risk consumers’ belief distortion outweighs the efficiency gains from pricing being better aligned with risk.

Relation to existing theoretical work

Our model is best related to the existing theoretical literature via its three key features: variation in willingness to pay across customers with identical risk (formulated here via differences in beliefs), a pooled insurance contract, and imperfect competition.

The theoretical literature on welfare in insurance markets with heterogeneous willingness to pay includes Finkelstein and McGarry (2006), Cohen and Einav (2007), Einav, Finkelstein and Cullen (2010), Spinnewijn (2013, 2017) and Handel, Kolstad and Spinnewijn (2019). Spinnewijn (2013) considers heterogeneous beliefs and includes a monopoly insurer, but applies this to insurers selling a menu of screening contracts. Spinnewijn (2017) uses additive heterogeneity in willingness to pay, motivated by heterogeneous beliefs, frictions and bounded rationality and applies this to pooled contracts, but uses a competitive market structure and does not consider the impact of heterogeneity on premiums.

Stiglitz (1977) studied a monopoly insurer with a range of risk types but did not incorporate heterogeneity in willingness to pay and focuses on screening contracts. Mahoney and Weyl (2017) provide a general, reduced-form framework for imperfect competition in pooled insurance markets with adverse selection, but do not derive a belief-based microfoundation for demand heterogeneity. The specific pricing and welfare results that we present here turn out to depend crucially on all three of the key features.

Roadmap

Section 2 sets out the baseline framework of agents with CRRA utility and fully “interior” demand. It derives the “neutrality” finding of premiums not depending on the extent of risk heterogeneity right up to the point where the market collapses. Section 3 alters the model so that belief deviations scale proportionally with objective risk, derives the result that monopoly premiums fall as adverse selection increases and compares this with outcomes under perfect competition and Cournot oligopoly. Section 4 presents a model where both objective risks and beliefs are uniformly distributed, while Section 5 considers both as truncated Normals. Section 6 discusses coinsurance. Section 7 analyzes the welfare consequences of imposing pooled pricing, comparing outcomes under pooled and risk-based pricing across the two specifications and market structures.

2. Common Belief Heterogeneity Across Risk Types

This section presents a model in which the extent of belief heterogeneity is unrelated to the level of objective risk.

2.1. Setup

We normalize initial wealth to one. Customers have a probability of incurring a loss of size $L \in (0, 1)$. There is a continuum of risk groups, indexed by i . Group i 's objective probability of loss is p_i , with population mean $E[p_i] = p$. Within group i , individual consumers face the same objective risk p_i but differ in their subjective beliefs about it.

There is a monopoly insurer who offers a single full insurance contract: a contract that pays the entire loss L if it occurs, in exchange for a premium R paid upfront. Consumers have CRRA utility in wealth

$$u(w) = \frac{w^{1-\rho}}{1-\rho} \quad (1)$$

for $\rho \neq 1$, with logarithmic utility $u(w) = \log(w)$ when $\rho = 1$. A consumer in group i with subjective belief $\tilde{p}$ about the loss probability has expected utility from remaining uninsured of

$$E[u_{unins}] = (1 - \tilde{p})u(1) + \tilde{p}u(1 - L) \quad (2)$$

With full insurance at premium R , certain utility is $u(1 - R)$. For realistically small premiums relative to wealth, we expand $u(1 - R)$ around $R = 0$ to give¹

$$u(1 - R) \approx u(1) - Ru'(1) \quad (3)$$

Under CRRA with wealth normalized to one, $u'(1) = 1$. The criteria for choosing to purchase insurance, $u(1 - R) \geq E[u_{unins}]$, becomes

$$u(1) - R \geq (1 - \tilde{p})u(1) + \tilde{p}u(1 - L) \quad (4)$$

Rearranging, this becomes

$$R \leq \tilde{p}[u(1) - u(1 - L)] \equiv \tilde{p}A \quad (5)$$

where

$$A \equiv u(1) - u(1 - L) = \frac{1 - (1 - L)^{1-\rho}}{1 - \rho} \quad (6)$$

¹The approximation error in the Taylor expansion $u(1 - R) \approx u(1) - Ru'(1)$ is of order $R^2 u''(1)/2$. Since the error is second order in R , it is small for any realistic premium. For example, with CRRA utility and risk aversion coefficient ρ , the relative error is of order $R^2 \rho/2$. At a premium of $R = 0.03$ — meaning three percent of wealth — and $\rho = 2$, this is approximately 0.09 percent, which is negligible for most purposes.

is the utility value of eliminating the loss, with $A = -\log(1 - L)$ under log utility. A consumer with subjective loss probability $\tilde{p}$ is therefore willing to pay up to $\tilde{p}A$ for the contract.

2.2. Belief Structure and Demand

Consider first the case where the distribution of the difference between beliefs and true loss probabilities is the same for all possible values of the true probability. This idiosyncratic component has a uniform distribution so individual j in risk group i has subjective loss probability

$$\tilde{p}_{ij} = p_i + \nu_j \quad \nu_j \sim U[-\gamma, \gamma] \quad (7)$$

where $\gamma > 0$ indexes the degree of idiosyncratic belief heterogeneity. They purchase the insurance if

$$(p_i + \nu_j)A \geq R \implies \nu_j \geq R/A - p_i \quad (8)$$

The fraction of consumers in risk group i who purchase is thus

$$D_i(R) = \frac{\gamma + p_i - R/A}{2\gamma} \quad (9)$$

truncated to $[0, 1]$, and the insurer's expected profit from risk group i is $(R - p_i L)D_i(R)$.

We assume that in the region of the profit-maximizing premium, demand for all risk groups lies between zero and one. This is essentially an assumption about the extent of belief heterogeneity relative to variations in true risk. It means that, at the optimal premium, even among the group with the lowest possible risk, someone is willing to purchase the insurance, while not everyone in the group with the highest risk is willing to purchase. These can be considered a useful baseline given the evidence cited above that variation in idiosyncratic willingness to pay appears to be large relative to actual risk heterogeneity. We later use an example with fully-specified risk heterogeneity to show that while this assumption can break down, it only tends to do so for unrealistic parameter combinations.

These assumptions imply the following result.

Proposition 1: *Let p_i have mean p and second moment $\sigma^2 \equiv E[p_i^2]$. Beliefs are given by equation 7 and assume the resulting demand $D_i(R)$, locally to the profit-maximizing value of R , lies in $[0, 1]$ for every risk group. Then the expected-profit-maximizing premium and profits are*

$$R^* = \frac{A\gamma + p(A + L)}{2} \quad (10)$$

$$\Pi^* = \frac{(R^*)^2}{2A\gamma} - \frac{Lp}{2} - \frac{L\sigma^2}{2\gamma} \quad (11)$$

and, letting $z \equiv A\gamma + p(A - L)$, profits are zero when

$$\sigma_0^2 = p^2 + \frac{z^2}{4AL} \quad (12)$$

and negative for $\sigma^2 > \sigma_0^2$: beyond this point, the pooled insurance contract is not profitable for a monopolist.

Proof: Averaging the demand function over the population, expected profit is

$$\Pi(R) = E \left[(R - p_i L) \left(\frac{\gamma + p_i - R/A}{2\gamma} \right) \right] \quad (13)$$

Expanding the term inside the expectation as a polynomial in p_i , we get

$$(R - p_i L)(\gamma + p_i - R/A) = -Lp_i^2 + \left(R + \frac{LR}{A} - L\gamma \right) p_i + \left(R\gamma - \frac{R^2}{A} \right) \quad (14)$$

Taking expectations over p_i and simplifying gives

$$\Pi(R) = \frac{1}{2\gamma} \left[R\gamma + Rp \left(1 + \frac{L}{A} \right) - \frac{R^2}{A} - L\gamma p - L\sigma^2 \right] \quad (15)$$

Differentiating with respect to R and solving gives

$$R^* = \frac{A\gamma + p(A + L)}{2} \quad (16)$$

Substituting R^* back into $\Pi(R)$ and collecting terms gives

$$\Pi^* = \frac{(R^*)^2}{2A\gamma} - \frac{pL}{2} - \frac{L\sigma^2}{2\gamma} \quad (17)$$

This is linear and strictly decreasing in σ^2 , so it has at most one zero. Solving $\Pi^* = 0$ for σ^2 gives $\sigma_0^2 = (R^*)^2/(AL) - \gamma p$, which simplifies directly to $p^2 + z^2/(4AL)$. ■

The model thus retains the classic market collapse result but it does not feature the gradual increase in premiums. The optimal premium depends on the average loss probability p but does not depend on any other moment of risk distribution. The adverse selection problem acts as subtraction from profits and does not affect pricing but, once this subtraction from profits gets large enough, the market collapses.

The mechanism is easiest seen in a simple two-type version of the model, with a high-risk group, $p_H = p(1 + \kappa)$ and a low-risk group, $p_L = p(1 - \kappa)$, of equal size. In this case, maximized profit is

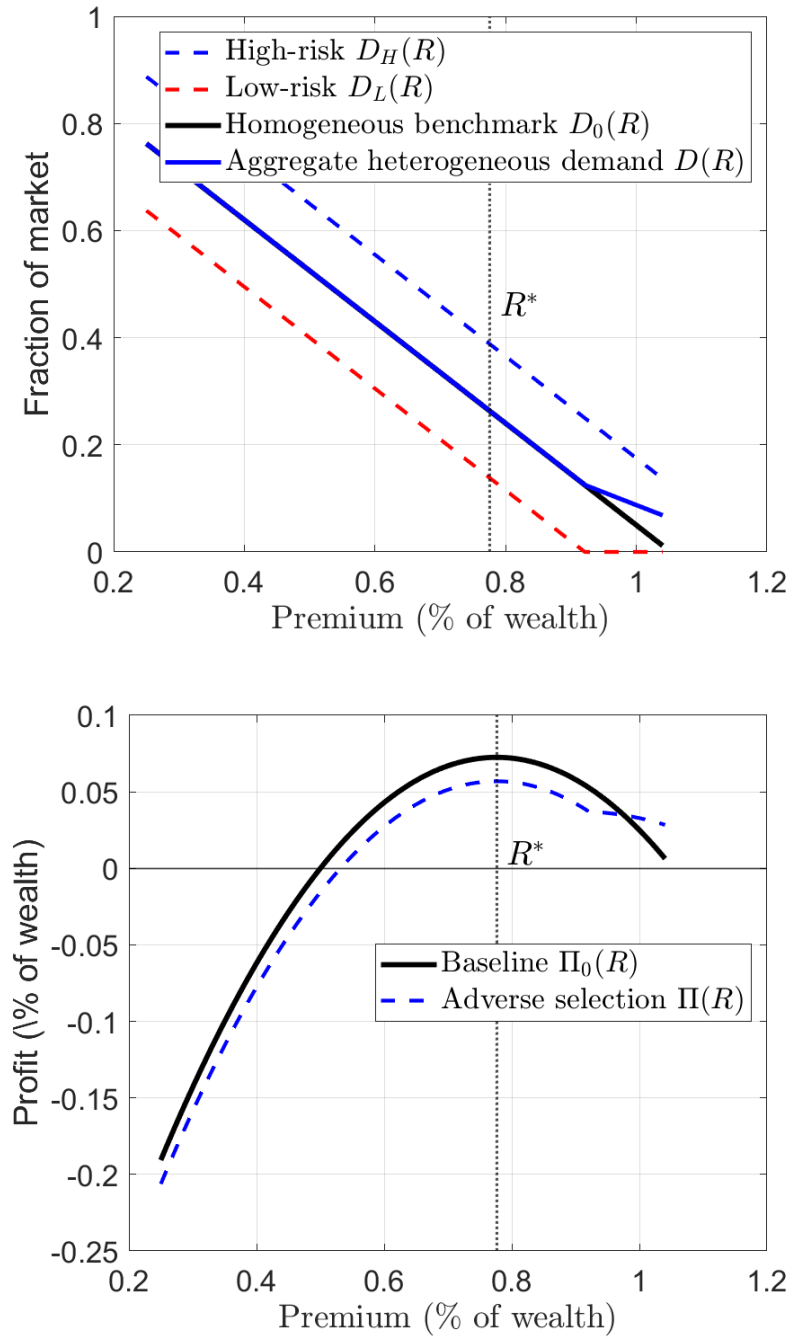
$$\Pi^*(\kappa) = \frac{(R^*)^2}{2A\gamma} - \frac{pL}{2} - \frac{p^2(1 + \kappa^2)L}{2\gamma} \quad (18)$$

which is strictly decreasing in κ : heterogeneity erodes profit through the adverse selection cost $Lp^2\kappa^2/(2\gamma)$, but leaves the price completely unaffected.

Figure 1 shows the demand and profit curves for $p = 0.10$, $L = 0.05$, $\rho = 2.0$ and $\kappa = 0.25$. Defining $D_0(R)$ as the baseline level of demand when there is only one risk type with a probability p of loss, the high-risk group's demand curve $D_H(R)$ lies above the baseline and the low-risk group's $D_L(R)$ lies below it. However, the average of D_H and D_L coincides exactly with D_0 at every premium up to and beyond the profit-maximizing level: the two demand curves are symmetric deviations around the homogeneous baseline, so the aggregate demand facing the monopolist is identical to what it would face with a homogeneous population at the mean risk p . Optimal pricing depends on the elasticity of this demand curve and so it is the same.

The graph shows that there is a point above the optimal premium where demand for the lower risk groups goes to zero and we see kinks in the aggregated demand and profit curves. These parameter values fit with the assumption of fully interior demands for both groups at the profit-maximizing premium, but this would not necessarily hold for higher values of κ .

Figure 1: Two-type model illustration with common belief heterogeneity for all risk groups
 $(p = 0.10, L = 0.05, \rho = 2.0, \kappa = 0.25)$.



3. Belief Heterogeneity Scaled With True Risk

Common additive belief heterogeneity for each value of p_i is simple and produces clean results but it is not very realistic. Suppose the heterogeneity parameter is $\gamma = 0.01$. This formulation means the variation in beliefs when $p_i = 0.02$ ranges from $p_{ij} = 0.01$ to $p_{ij} = 0.03$ while the same variation when $p_i = 0.10$ goes from $p_{ij} = 0.09$ to $p_{ij} = 0.11$. With a 2% loss probability, the most pessimistic person believes the risk is three times higher than the most optimistic person's belief, while a 10% probability of risk sees only a 22% difference between the most pessimistic and optimistic person. This seems inconsistent with the evidence from Sydnor (2010) and Abito and Salant (2019) that overestimation of loss risks scales with objective risk. Here we present a specification in which belief heterogeneity scales with the actual risk level.

3.1. Belief Structure and Demand

We assume now that individual j in risk group i has subjective loss probability

$$\tilde{p}_{ij} = p_i(1 + \nu_j) \quad \nu_j \sim U[\ell, h] \quad (19)$$

where ν_j is independent of p_i . A consumer cannot believe their risk is less than zero, so underestimation is bounded, $\ell \geq -1$, while overestimation is not. We work with $\ell = -1$, the natural floor at which the most skeptical consumer believes they face no risk whatsoever ($\tilde{p}_{ij} = 0$); h is left free and can exceed 1, so the most pessimistic consumer's belief is not capped at twice the true risk, consistent with the evidence of systematic over-pessimism among those who purchase insurance products.

Consumer (i, j) purchases insurance if

$$p_i(1 + \nu_j)A \geq R \quad (20)$$

giving demand

$$D(p_i, R) = 1 - \frac{R}{Ap_i(h - \ell)} = 1 - \frac{R}{Ap_i(1 + h)} \quad (21)$$

truncated to $[0, 1]$, where the last equality uses $\ell = -1$. In this case, demand is linear in $1/p_i$ rather than in p_i itself, which drives the key result below.

Proposition 2: *Let p_i have mean p and finite $q \equiv E[1/p_i]$. Beliefs are given by equation 19 and assume the resulting demand $D(p_i, R)$ lies in $[0, 1]$ for every risk group in the population at the profit-maximizing premium. Then*

$$R^* = \frac{A(1 + h) + L}{2q} \quad (22)$$

and any mean-preserving increase in the spread of p_i strictly increases q and hence strictly decreases R^ . Let*

$\beta \equiv A(1 + h) + L$. Then

$$\Pi^*(q) = \frac{\beta^2}{4A(1 + h)q} - pL \quad (23)$$

Π^* is strictly decreasing in q and equals zero exactly at

$$q_0 \equiv \frac{1}{p} + \frac{[A(1 + h) - L]^2}{4ALp(1 + h)} \quad (24)$$

and $\Pi^* < 0$ for $q > q_0$: beyond this point, the pooled insurance contract is not profitable for a monopolist. The aggregate quantity sold at R^* ,

$$Q^*(q) = 1 - \frac{R^*q}{A(1 + h)} = \frac{A(1 + h) - L}{2A(1 + h)} \quad (25)$$

does not depend on q : however risk is distributed across the population, the monopolist sells exactly the same total quantity, adjusting only the price at which it does so.

Proof: The insurer's expected profit is

$$\Pi(R) = E \left[(R - Lp_i) \left(1 - \frac{R}{Ap_i(1 + h)} \right) \right] \quad (26)$$

Taking expectations gives

$$\Pi(R) = R - Lp - \frac{R^2q}{A(1 + h)} + \frac{LR}{A(1 + h)} \quad (27)$$

Differentiating with respect to R and solving gives

$$R^* = \frac{A(1 + h) + L}{2q} \quad (28)$$

R^* is strictly decreasing in q for any $q > 0$. Since $1/p_i$ is a strictly convex function of p_i , Jensen's inequality implies that any mean-preserving spread of p_i — holding $E[p_i] = p$ fixed while increasing dispersion — strictly increases $q = E[1/p_i]$. Greater risk heterogeneity therefore strictly lowers the profit-maximizing premium.

Substituting $R^* = \beta/(2q)$ into $\Pi(R)$ and simplifying gives $\Pi^*(q)$ as stated; it is a strictly decreasing function of q for $q > 0$, so it has at most one zero. Solving $\Pi^* = 0$ for q gives $\beta^2 = 4A(1 + h)Lpq$, i.e. $q_0 = \beta^2/[4ALp(1 + h)]$, which simplifies to $1/p + [A(1 + h) - L]^2/[4ALp(1 + h)]$.

Aggregate quantity at price R is $Q(R) = 1 - Rq/[A(1 + h)]$, from averaging $D(p_i, R)$ across the population. Substituting $R^* = \beta/(2q)$ directly:

$$Q^*(q) = 1 - \frac{1}{A(1 + h)} \frac{\beta}{2q} q = 1 - \frac{\beta}{2A(1 + h)} = \frac{A(1 + h) - L}{2A(1 + h)} \quad (29)$$

The q inside R^* and the q multiplying it in $Q(R)$ cancel exactly, leaving no q in the result. ■

The model again retains the standard intuition that sufficient risk heterogeneity can eventually cause the market to collapse. But prior to market collapse, the effect on pricing is the opposite of the usual intuition. As risk heterogeneity increases, the monopoly insurer reduces premiums.

Proposition 2 places no restriction on how risk is distributed across the population, other than the requirement that $D(p_i, R) \in [0, 1]$. If there is only one risk type, $q = 1/p$ and the monopolist charges $p\beta/2$, selling Q^* . If instead there are several observably different risk groups and the monopolist can set a separate price for each, its profit is additively separable across groups — neither the demand nor the cost of serving group i depends on the price charged to any other group — so pricing group i optimally is exactly the single-group problem with $q = 1/p_i$: it charges $p_i\beta/2$ and again sells Q^* , the same quantity as every other group. Full price discrimination across risk types and pooling everyone into a single price therefore sell exactly the same aggregate quantity; only the price charged, and who ends up buying it, differs between the two.

3.2. Two-Type Intuition

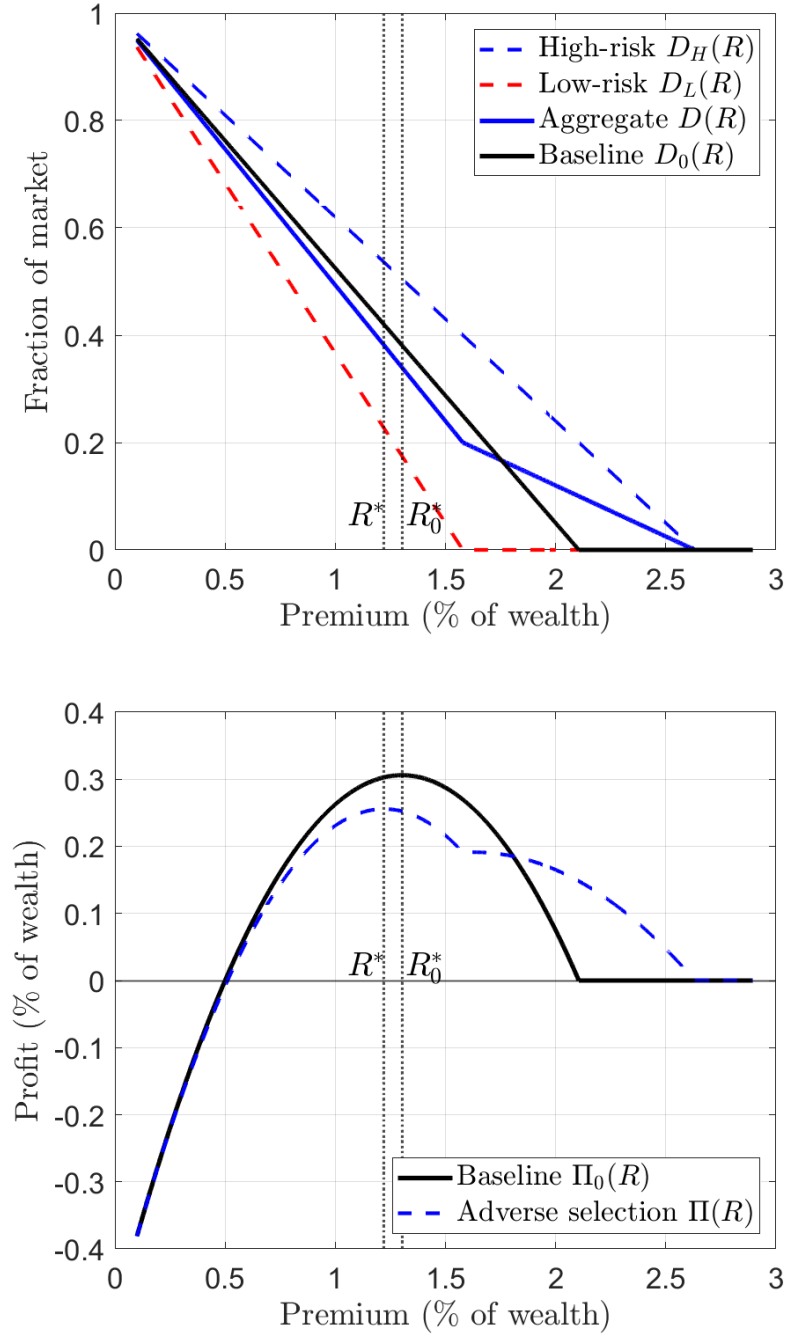
The intuition for the behavior of premiums described in Proposition 2 is again easiest explained with a simple two-type version of the model. Figure 2 shows demand and profit curves for $p = 0.10$, $L = 0.05$, $\rho = 2.0$ and $\kappa = 0.25$.

The low- and high-risk demand curves are no longer parallel. High-risk consumers have wider belief distributions in absolute terms — their willingness-to-pay $p_H(1 + \nu_j)A$ ranges over a wider dollar interval than the low-risk group's — making their demand curve less steep than the homogeneous baseline. Low-risk consumers, in contrast, are more price sensitive and this is the dominant effect on the aggregate demand curve's slope. The slope of the demand curves scale with $1/p_i$ and the average of $1/p(1 - \kappa)$ and $1/p(1 + \kappa)$ is larger than $1/p$ due to Jensen's inequality.

As a result, the aggregate demand curve $D(R)$ lies below the homogeneous baseline $D_0(R)$ at any given premium: the pool is, in aggregate, more elastic than a homogeneous population at the mean risk p . This shifts the profit-maximizing premium downward: adverse selection makes the marginal consumer more price-sensitive relative to the average, so the monopolist gains more from cutting the price to retain them than from raising it. The chart also shows the optimal premium R_0^* from the homogeneous case where all customers have loss risk of p .

The charts again show that as premiums rise above the profit-maximizing level, low-risk demand goes to zero and there is a kink in the profit function. In this case, the global optimal profit level is prior to this kink but, as we discuss later, for sufficiently high risk heterogeneity (most likely, unrealistically high), this “truncated” case can be the global optimal solution for the insurer.

Figure 2: Two-type model illustration with belief heterogeneity scaling with risk ($p = 0.10$, $L = 0.05$, $\rho = 2.0$, $\kappa = 0.25$). R_0^* is the premium if there was one type with loss rate equal to the average of the two-type case.



3.3. Other Market Structures

To further explore the role of market structures in driving these results, here we consider a zero-profit competitive market and a Cournot oligopoly.

Proposition 3. (Proof in Appendix A.1). *Let q , q_0 and β be defined as in Proposition 2 and retain the same assumptions about beliefs and demands, then if $q \leq q_0$, the zero-profit competitive premium is*

$$R^{PC} = \frac{\beta - \sqrt{\beta^2 - 4ALp(1+h)q}}{2q} \quad (30)$$

R^{PC} equals the actuarially fair premium pL when $q = 1/p$ and is strictly increasing in q throughout. If $q > q_0$, no premium allows insurers to break even while serving every risk group.

Proposition 4. (Proof in Appendix A.2). *Let q , q_0 and β be defined as in Proposition 2 and retain the same assumptions about beliefs and demands. Consider n symmetric firms that share aggregate claims cost in proportion to output share. A symmetric Cournot equilibrium with non-negative profit exists if and only if $q \leq q_0$ for every n . When it exists, the total quantity of insurance sold is*

$$Q_n^* = \frac{n[A(1+h) - L] + \sqrt{n^2[A(1+h) - L]^2 - 4A(1+h)(n+1)L(n-1)(pq-1)}}{2A(1+h)(n+1)} \quad (31)$$

and premiums are

$$R_n^* = A(1+h)(1 - Q_n^*)/q \quad (32)$$

As $q \rightarrow q_0$, $Q_n^* \rightarrow Q^* = [A(1+h) - L]/[2A(1+h)]$ for every n and, for any fixed $q \leq q_0$, $R_n^* \rightarrow R^{PC}$ as $n \rightarrow \infty$.

These propositions show the crucial role played by market structure in deriving the result that increasing risk heterogeneity drives lower premiums. Competitive pricing implies the traditional result of gradually higher premiums followed by market collapse holds in the presence of heterogeneous beliefs. The Cournot results show that the pricing will generally be in between competitive and monopoly, so whether premiums rise or fall with mean-preserving spreads of the risk distribution will depend on the extent of competition in the market.

4. Uniformly Distributed Loss Probabilities

We have thus far only considered cases where there is an interior solution to the demand problem for all risk groups. This was an implicit assumption about belief heterogeneity being sufficiently large relative to risk heterogeneity, and it allowed for very general results about the impact of risk heterogeneity on premiums and volumes. Here, we use an explicit specification for risk heterogeneity. We assume risk probabilities are uniformly distributed. Specifically, we assume

$$p_i = p(1 + \delta_i) \quad \delta_i \sim U[-\kappa, \kappa] \quad (33)$$

and use this formulation of risk to derive solutions for the optimal premium when interior solutions do not always hold.

4.1. Monopoly

We first show how monopoly pricing changes with this new specification of risk heterogeneity. The original premium formula is correct as long as there is an “interior” solution with demand for each group between one and zero. But once this condition does not hold, both pricing and market collapse results change.

Proposition 5. (Proof in Appendix A.3). *Let $p_i = p(1 + \delta_i)$, $\delta_i \sim U[-\kappa, \kappa]$, and let $\nu_j \sim U[-1, h]$, so demand is $D(\delta_i, x) = \max\{0, \min\{1, 1 - x/[p_i(1 + h)]\}\}$ where $x \equiv R/A$. Let $q(\kappa) \equiv E[1/p_i] = \ln[(1 + \kappa)/(1 - \kappa)]/(2\kappa p)$, and let κ^* be the unique value of $\kappa \in (0, 1)$ solving*

$$\kappa^* [A(1 + h) + L] = A(1 - \kappa^*)(1 + h) \ln\left(\frac{1 + \kappa^*}{1 - \kappa^*}\right) \quad (34)$$

The profit-maximizing premium is

$$R^*(\kappa) = \begin{cases} \frac{A(1 + h) + L}{2q(\kappa)}, & \kappa \leq \kappa^* \\ A(1 + h)c_M p(1 + \kappa), & \kappa > \kappa^* \end{cases} \quad (35)$$

where $B \equiv [A(1 + h) + L]/[2A(1 + h)]$ and c_M is the unique solution in $(0, 1)$ to $B(1 - c_M) = -c_M \ln c_M$, depending only on B — hence only on A , L , and h — not on κ . This function is continuous and so does not have a jump at κ^ .*

This proposition has a number of implications, proved formally as corollaries in appendix A.4. For a fixed set of parameters, p , L , ρ and h , premiums fall as κ rises up to κ^* and then rise linearly with it. Economically, c_M is the marginal (lowest-risk) buyer’s risk as a fraction of the population’s highest risk, $\bar{p} = c_M p(1 + \kappa)$. The proof of Proposition 5 shows that this fraction is fixed by B alone, and is independent of κ . This means the truncated-regime premium $A(1 + h)c_M p(1 + \kappa)$ is a constant

markup applied to the population's risk ceiling, which is why the premium rises linearly in κ once truncation begins.

Note that the κ^* here is distinct from our earlier results about market collapse. From equation 24 for the zero profit point for monopolists, you can obtain an implicit equation for the level of κ consistent with zero profits under the interior-regime premium. However, that value is irrelevant here, since κ^* is below it. By the time risk heterogeneity has reached that value of κ , lower-risk groups are already fully dropping out, rendering the previous aggregate demand formula (and its collapse prediction) incorrect. Above κ^* , the optimal premium now involves both the implications of higher premiums for interior demand functions but also a Leibniz rule element in which the lower limit of the integral also changes with premiums. The result is a kink in the optimal premium, so that premiums shift from falling to rising linearly with extent of risk heterogeneity.

In relation to volumes, Proposition 2 already showed that quantity in the interior regime is independent of κ . However, volumes fall once truncation excludes the lowest-risk types,

$$Q_{II}^*(\kappa) = \frac{(1 + \kappa)[(1 - c_M) + c_M \ln c_M]}{2\kappa} \quad (36)$$

Since $\nu_j \sim U[-1, h]$ is only bounded below, a high-risk individual with $p_i = p(1 + \kappa)$ and a highly pessimistic belief draw $\nu_j = h$ perceives a loss probability of $p(1 + \kappa)(1 + h)$, which exceeds one once κ and h are both large enough. To keep every displayed parameterization economically meaningful, κ is capped at 0.8 in every figure in this section, satisfying $p(1 + \kappa)(1 + h) \leq 1$ for every h shown; this only binds for h close to 4.5 and is well above the range of κ relevant to the results discussed below.

Figures 3 and 4 plot $R^*(\kappa)$ and $\Pi^*(\kappa)$, with $p = 0.10$, $L = 0.05$, $\rho = 2.0$, $\ell = -1$, and h ranging from 1 to 4.5 in steps of 0.5. The solid portion of each curve shows outcomes when all groups' demands are between zero and one and the dashed portion is the truncation branch.

Figure 3 shows premiums falling from their homogeneous-risk benchmark, reaching a minimum at $\kappa^*(h)$, and then rising linearly through the truncated regime. Both branches shift upward and $\kappa^*(h)$ moves right as h increases. In Figure 4, by contrast, $\Pi^*(\kappa)$ passes smoothly through $\kappa^*(h)$ with no kink. Below $\kappa^*(h)$, maximized profits fall in a concave fashion as κ rises but this decline then switches to being convex, so the market never collapses. The move to the truncated solution, and thus rising premia with increased risk heterogeneity, starts at about $\kappa = 0.28$ for the lowest value of h . This already implies a significant amount of risk heterogeneity—the riskiest group has a loss probability almost twice the lowest risk group—and this threshold level rises for higher levels of h , which appear more consistent with the evidence cited above.

Another way to summarize the same model is in terms of the profit rate on premiums actually accepted, $\Pi^*(\kappa)/[R^*(\kappa)Q^*(\kappa)]$ — profit per dollar of premium revenue collected. Figure 5 plots this profit rate for the same parameters and the same range of h . This rises in h and falls in κ up to the

truncation point, after which it is flat.

At $h = 1$, $\nu_j \sim U[-1, 1]$ has mean zero, so the average consumer’s perceived risk equals the true risk exactly, yet the profit rate at $\kappa = 0$ already ranges from 29% to 36%. The mechanism in this case is not biased beliefs — customers are on average correct in their risk assessments — but disagreement: consumers underestimate and overestimate their risk in offsetting amounts but the product is priced for those who overestimate so the price can be well above the actuarially fair rate. A similar result has been provided by Hegarty and Whelan (2026) in the context of sports betting: bookmakers’ profits rise with disagreement among bettors about probabilities, even if the bettors are on average correct in their beliefs. For $h > 1$, the population’s average belief is biased upwards, further reinforcing higher premiums. Profit rates at $h = 4.5$ range from 65–71%.

These magnitudes are broadly consistent with the empirical evidence on high profit rates for extended warranties in Abito and Salant (2019). They find that probability distortion is the dominant driver of both demand and profit in that market, and report back-of-envelope margins of 62–73% on TV warranties (their Table 2) alongside a structurally estimated retailer margin of roughly 46%. Our h parameter plays the analogous role to their probability-distortion estimates: in both cases, a larger upward distortion of perceived risk is associated with a substantially higher profit share, holding the underlying cost and market structure fixed. These high profit rates may reflect people having less accurate beliefs about unusual events like appliances breaking than about their own health or the chance of their car breaking down.

The results here temper the new element of our previous results (the flat premium) while undermining the traditional element (the market collapse prediction). In relation to pricing, our previous results showed premiums falling as risk heterogeneity increased, but incorporating some risk groups choosing not to purchase any insurance implies that premiums will rise once adverse selection problems become severe enough. But we also find that our earlier results about markets collapsing as adverse selection problems worsen may not hold. Under this model of risk heterogeneity, the market instead moves to dropping low-risk customers. Our numerical results suggest, however, that the new regime implied here—with low-risk customers dropping out altogether and premiums rising with risk heterogeneity—requires high levels of risk heterogeneity relative to belief heterogeneity. This may mean the “interior” regime results are the more empirically relevant ones.

Figure 3: The optimal premium R^* as a function of risk heterogeneity κ , for $h \in \{1, 1.5, \dots, 4.5\}$ (with $p = 0.10, L = 0.05, \rho = 2.0, \ell = -1$).

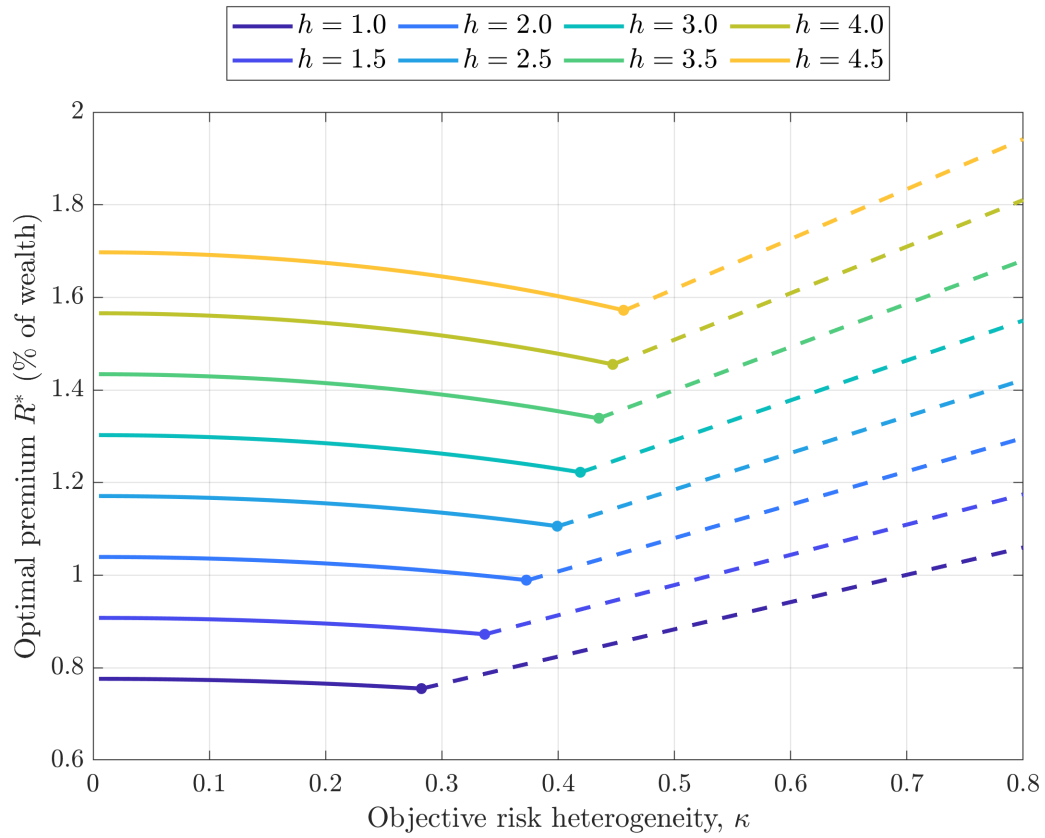


Figure 4: Maximized profit Π^* as a function of risk heterogeneity κ , for the same values of h and parameters as Figure 3.

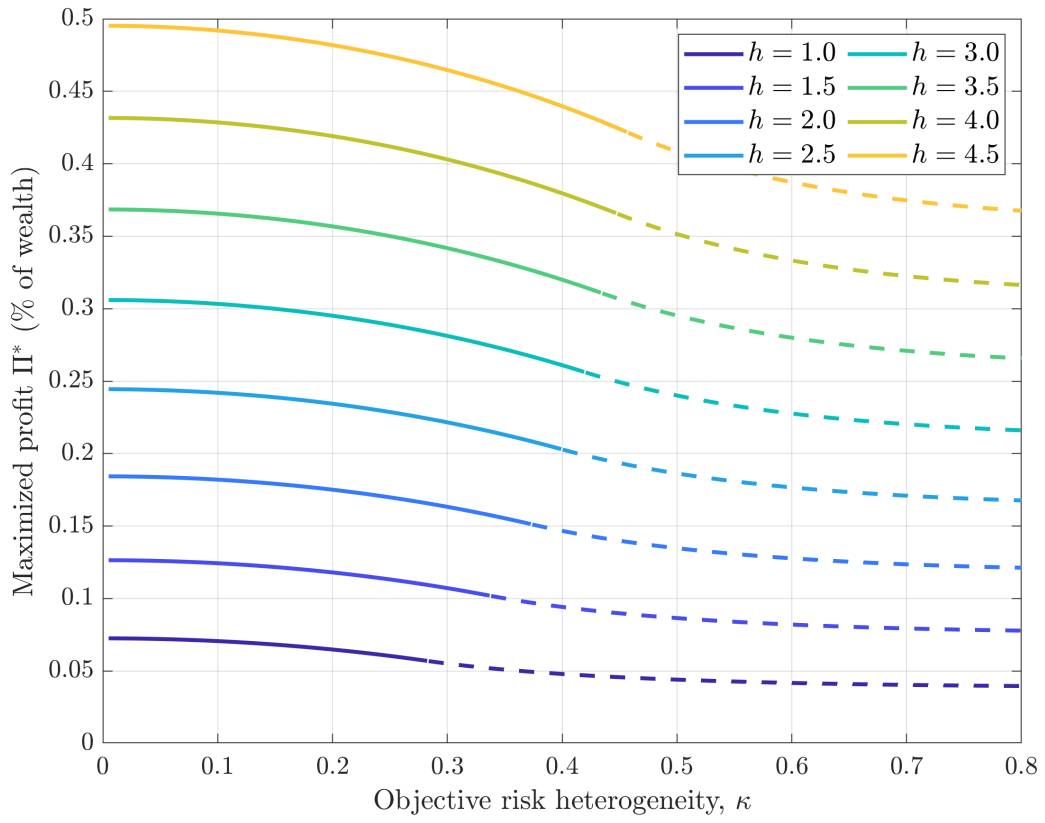
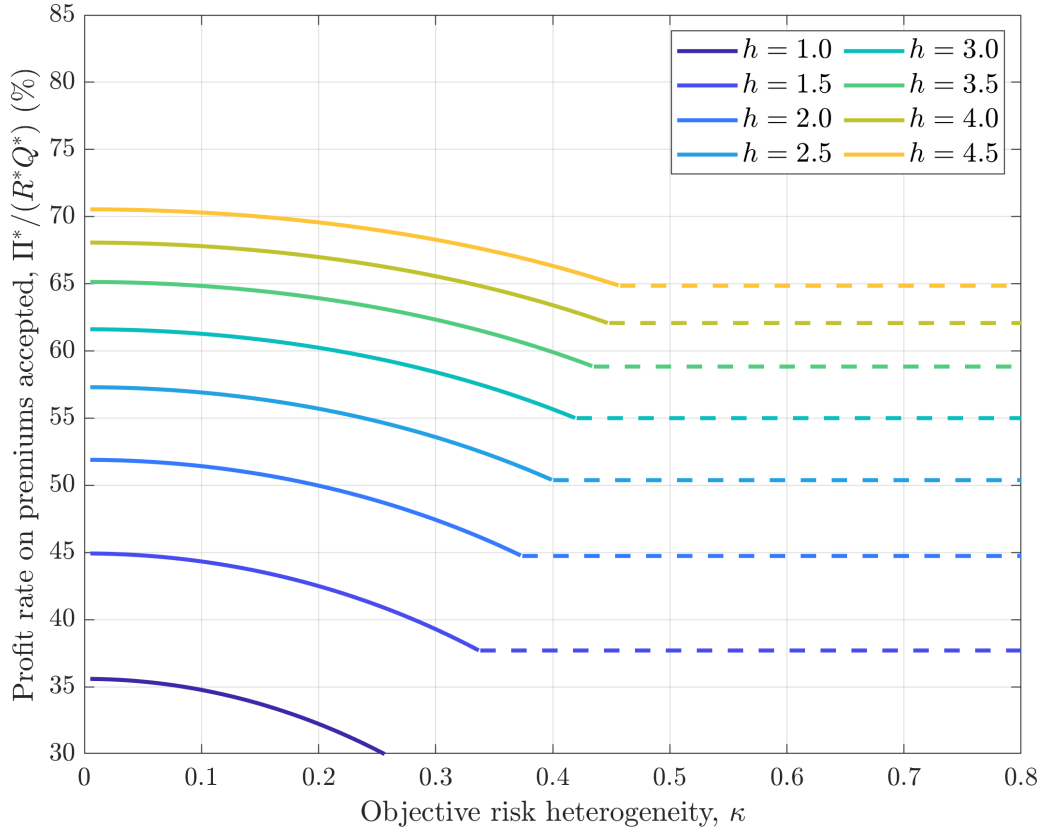


Figure 5: Profit rate on premiums accepted, $\Pi^*/(R^*Q^*)$, as a function of risk heterogeneity κ , for the same values of h and parameters as Figure 3.



4.2. Competition

The role of market structure in generating our findings can again be emphasized by deriving the competitive equilibrium for this model of risk heterogeneity.

Proposition 6. (Proof in Appendix A.5). *Let the distributions of risk probabilities and beliefs be as assumed in Proposition 5 and let β and $q(\kappa)$ be as defined earlier. Let $\kappa^{trunc,C}$ be the value of κ at which the interior zero-profit premium $R^{PC}(\kappa)$ as defined in equation 30 equals $R_{exit}(\kappa) \equiv Ap(1 - \kappa)(1 + h)$, the price at which the lowest-risk type's demand reaches exactly zero. For $\kappa \leq \kappa^{trunc,C}$, the equilibrium premium $R^{PC}(\kappa)$ is the interior formula of Proposition 3. For $\kappa > \kappa^{trunc,C}$,*

$$R^{PC}(\kappa) = A(1 + h)c_C p(1 + \kappa) \quad (37)$$

where $c_C \in (0, 1)$ solves

$$2c_C^2 \ln c_C - (2B + 1)c_C^2 + 4Bc_C - 2B + 1 = 0, \quad B \equiv \frac{\beta}{2A(1 + h)} \quad (38)$$

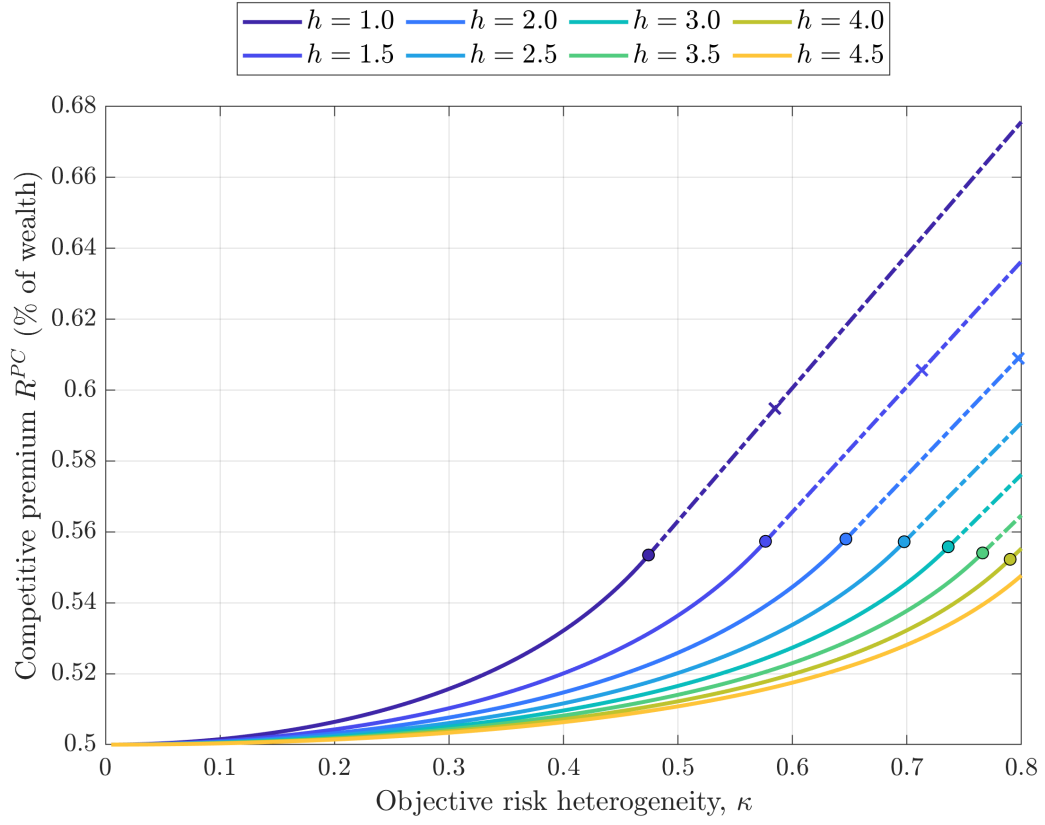
c_C depends only on B , not on κ . Provided $Ac_M(1 + h) > L$, where c_M solves the (different) equation $B(1 - c_M) = -c_M \ln c_M$ from Proposition 5's proof, the market never collapses: $R^{PC}(\kappa)$ exists for every $\kappa \in (0, 1)$, and the population share still served,

$$\text{Share}(\kappa) = \frac{(1 + \kappa)(1 - c_C)}{2\kappa} \quad (39)$$

converges to the strictly positive limit $1 - c_C$ as $\kappa \rightarrow 1$.

Note that condition $Ac_M(1 + h) > L$ implies that the premium exceeds the cost of insuring the riskiest, least profitable type in the population, $Lp(1 + \kappa)$. Margin per policy falls in risk, so if even that type is served profitably, every buyer is. Figure 6 plots $R^{PC}(\kappa)$ for the same values of h used throughout this paper. Unlike the monopoly premium's V-shape, each curve rises monotonically, continuing smoothly past $\kappa^{trunc,C}$ rather than terminating there.

Figure 6: The competitive (zero-profit) premium R^{PC} as a function of risk heterogeneity κ , for the same values of h and parameters as Figure 3. Unlike the monopoly premium, R^{PC} rises monotonically and never collapses.



4.3. Cournot Oligopoly

The following establishes the Cournot solution for this specification of risk heterogeneity.

Proposition 7. (Proof in Appendix A.6). Let $p_i = p(1 + \delta_i)$, $\delta_i \sim U[-\kappa, \kappa]$, $\nu_j \sim U[-1, h]$, let $\beta \equiv A(1 + h) + L$, $B \equiv \beta/[2A(1 + h)]$, and consider n symmetric firms that share aggregate claims cost in proportion to output share. For κ below a threshold κ_n (no closed form; rises from κ^* at $n = 1$ to $\kappa^{trunc, C}$ as $n \rightarrow \infty$), $Q_n^*(\kappa)$ is the interior Cournot solution of Proposition 4, specialized to uniform δ_i . For $\kappa > \kappa_n$,

$$R_n^*(\kappa) = A(1 + h)c_n p(1 + \kappa) \quad (40)$$

where $c_n \in (0, 1)$ solves

$$n \ln(c_n) [2c_n^2 \ln c_n - (2B + 1)c_n^2 + 4Bc_n - 2B + 1] + \text{Rem}(c_n) = 0 \quad (41)$$

with

$$\text{Rem}(c) \equiv 4B(c - 1)^2 + 2B \ln(c)(1 - c^2) + 2c^2 \ln^2 c - \ln(c)(3c - 1)(c - 1) \quad (42)$$

c_n depends only on B and n , not on κ . At $n = 1$, c_n equals monopoly's own constant (Proposition 5, solving $B(1 - c_M) = -c_M \ln c_M$). As $n \rightarrow \infty$, the bracketed term must vanish, so $c_n \rightarrow c_C$, the competitive truncated constant of Proposition 6.

Figure 7 plots the premium under monopoly, Cournot with $n = 2, 3, 10$, and perfect competition together at a single value $h = 2$. The rankings follow the extent of market power, with premiums falling as the number of competitors rises. For Cournot, premiums initially fall with κ when there are a small number of competitors but switch to always rising as the number of competitors gets big enough.

Figure 8 shows the corresponding volumes – the fraction of customers that buy the insurance – for the same parameter values. This shows that variations in market power produce large deviations in the fraction buying insurance, with volumes more than twice as high under competition as under monopoly. Interestingly, however, while quantities typically decline as risk heterogeneity increases, none of these markets come close to collapsing, with all quantities remaining well above zero, even at the maximum level of adverse selection possible for this calibration, $\kappa = 1$.

In relation to which of the regimes is likely to correspond to reality – premiums falling as risk heterogeneity or rising with it — the lowest transition point is for the monopoly case, where the transition occurs at $\kappa \approx 0.37$. This is a situation in which the highest risk group is 2.19 times more risky than the lowest group, which is already a substantial amount of heterogeneity. These truncation points rise as competition increases. Calculations also show them rising as h increases. Overall, it seems possible that the regime with falling premiums is an empirically relevant one.

Figure 7: Premiums under monopoly, Cournot oligopoly ($n = 2, 3, 10$), and perfect competition, all at a single fixed $h = 2.0$ (with $p = 0.10$, $L = 0.05$, $\rho = 2.0$, $\ell = -1$), as functions of risk heterogeneity κ . Circles and the square mark each curve's terminal point.

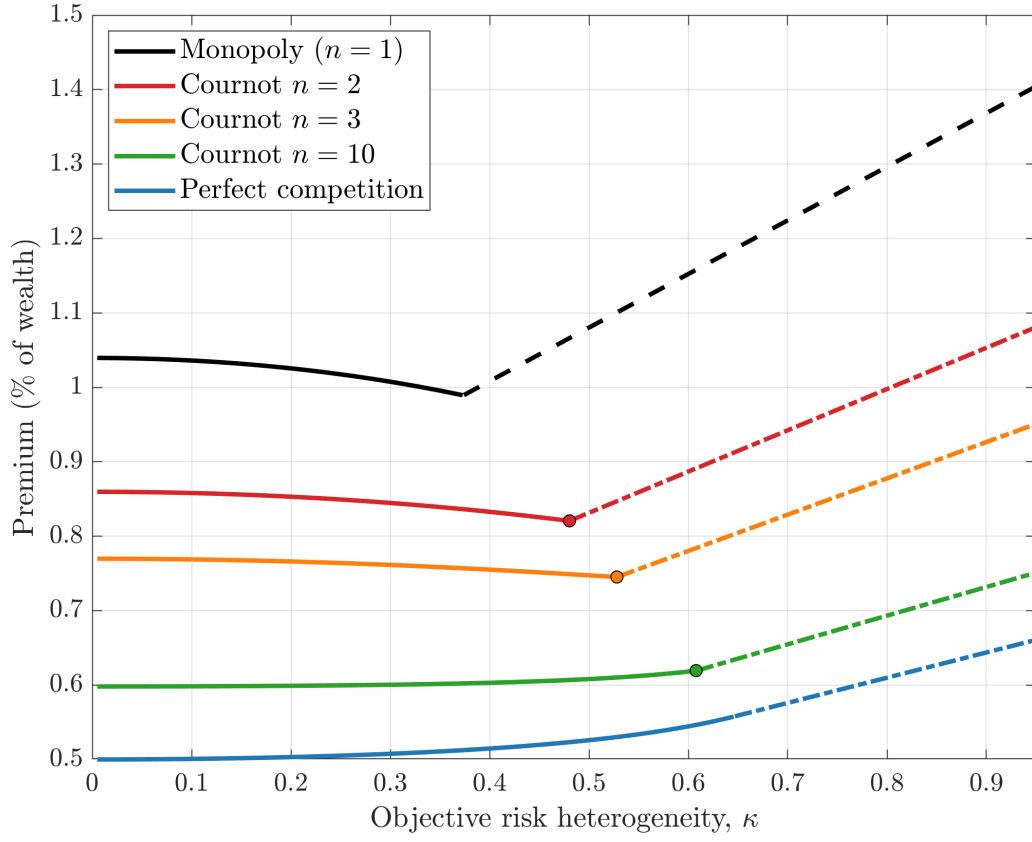
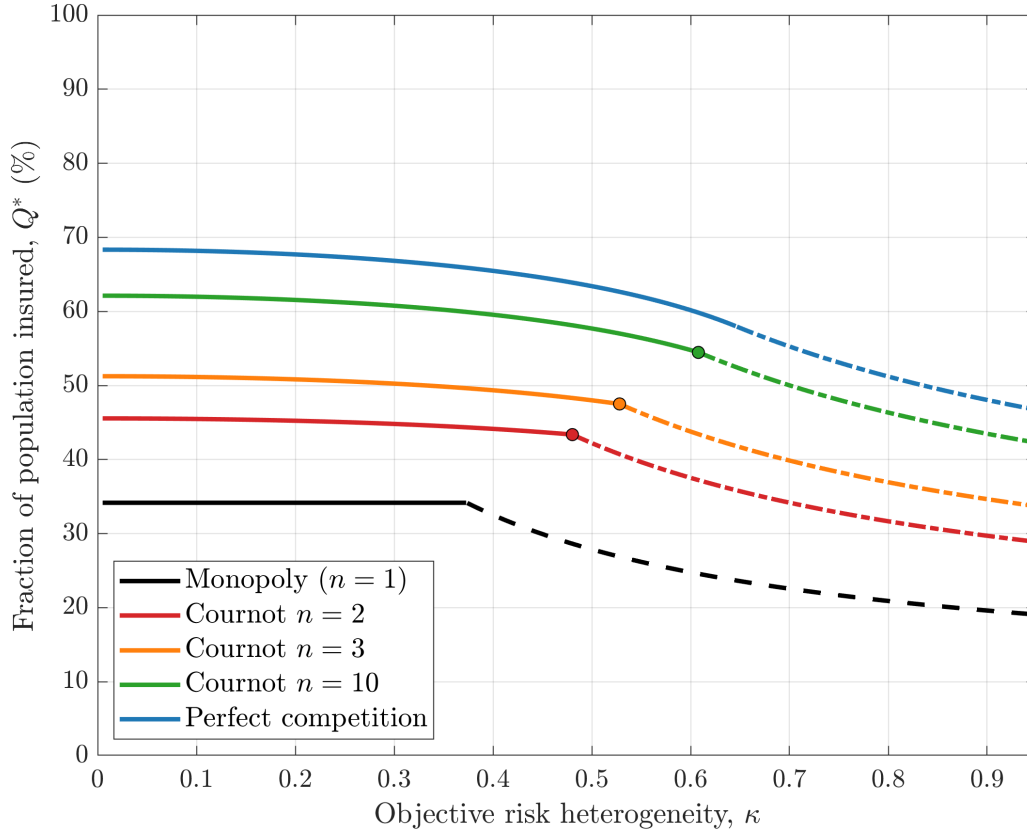


Figure 8: Fraction insured under monopoly, Cournot oligopoly ($n = 2, 3, 10$), and perfect competition, all at a single fixed $h = 2.0$ (with $p = 0.10$, $L = 0.05$, $\rho = 2.0$, $\ell = -1$), as functions of risk heterogeneity κ . Circles and the square mark each curve's terminal point.



5. Truncated Normal Distributions

Propositions 5, 6 and 7 assume both δ_i and ν_j are drawn from uniform distributions. The uniformity assumption for δ_i is not always relevant. As described above, within the interior regime of demand for all groups in $(0, 1)$, the formulas for premiums in Propositions 2 to 4 hold for any mean- p distribution of objective risk with finite $E[1/p_i]$, with uniformity entering only through the closed-form value of $q(\kappa)$ itself. However, uniformity of ν_j has been a key assumption all along, since it has determined the precise form of the demand function for each risk group.

To check how influential the uniformity assumption is, we also solved the monopoly version of the model numerically for a range of values of h and κ with lower-truncated Normal distributions for both p_i and ν_j that have the same mean and variance as the uniform distributions we have been using. The objective risk p_i was truncated at zero and the proportional belief parameter ν_j was truncated at -1, ruling out negative probability beliefs.² As in Section 4, κ is capped at 0.8 so that no individual's perceived loss probability, $p_i(1 + \nu_j)$, exceeds one for any of the values of h shown.

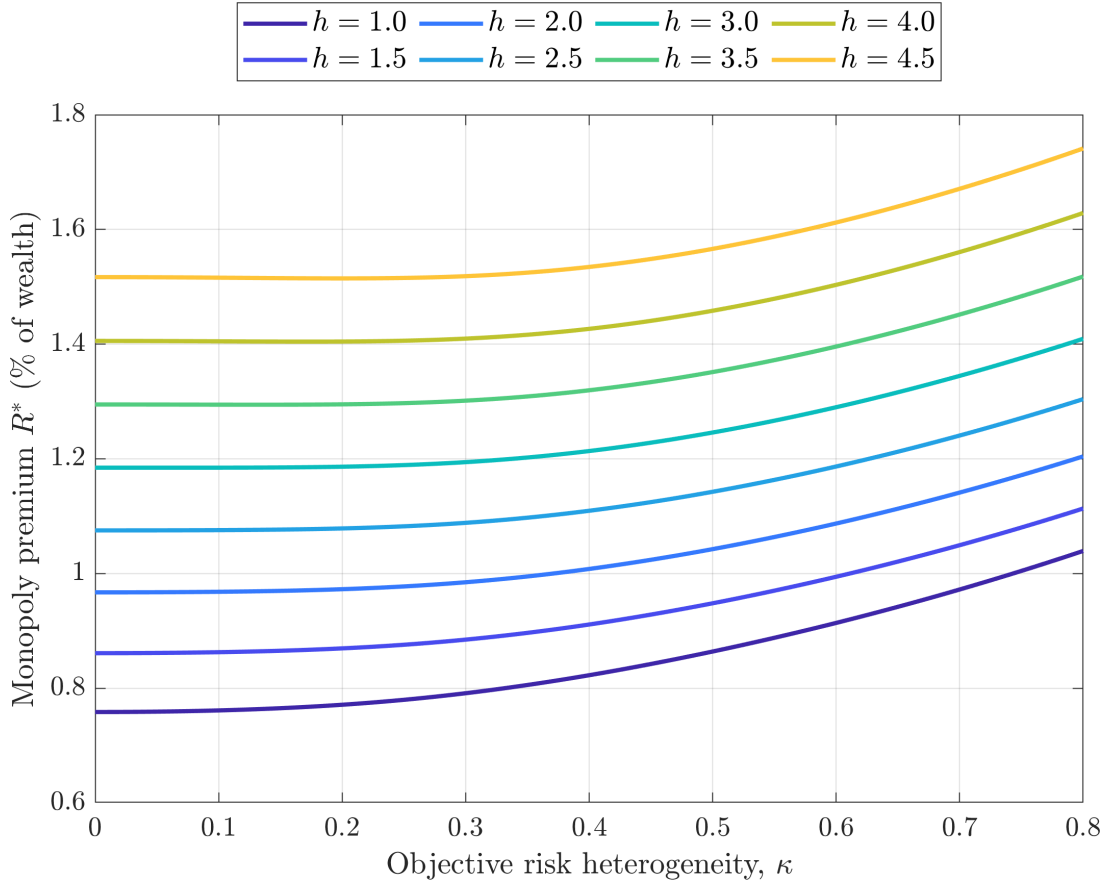
Figure 9 shows the profit-maximizing monopoly premiums for the same range of different specifications of both risk and belief heterogeneity. The truncated distributions are defined by means and variances rather than the single parameter values h and κ . What we show however, is the premiums obtained from the truncated Normal distributions that correspond to the same means and variances as the displayed values of these parameters generate in the uniform case.

The figure shows premiums still initially decline as κ rises from low values, though the effect is much weaker than with uniform distributions. They then transition to rising with κ in a more smooth fashion than under uniform distributions: the sharp kink shown before was a consequence of the uniform distribution. Once premiums begin to rise, volumes behave similarly to the uniform distribution case.

Overall, these results are compatible with the earlier findings. There is a realistic range of values for which worsening adverse selection does not raise pooled premiums under monopoly. And, rather than collapse, the market responds to higher rates of risk heterogeneity by raising premiums in a smooth fashion with volumes gradually declining but never going to zero.

²Section B has a full description of the numerical methods used.

Figure 9: The optimal premium R^* with truncated Normal distributions for p_i and ν_j with means and variances equal to uniform distributions with the displayed values of h and κ (with $p = 0.10$, $L = 0.05$, $\rho = 2.0$, $\ell = -1$).



6. Coinsurance

So far, we have assumed the insurer only offers full coverage. However, a monopoly insurer might instead prefer to provide only partial coverage, for instance via requiring a deductible. Here, we extend the proportional beliefs model of Section 3 to let the monopolist choose a coinsurance rate jointly with the premium, so that only a fraction ω of the loss is covered, with the customer incurring the rest.

Incorporating the insurer being able to pick both a premium and a coinsurance rate into the model in Section 4 does not produce an analytical solution for the profit-maximizing policy. Instead, we solved for the optimal policy in two steps. First, we fixed a specific coinsurance rate ω and obtain an analytical solution for the profit-maximizing premium for this rate. Then we used numerical methods to search over the range of possible values for ω to find the value associated with the highest profits.

The fixed coinsurance rate problem generalizes two objects used in Section 4. A is replaced by $A_\omega \equiv u(1 - (1 - \omega)L) - u(1 - L)$ and L is replaced by $L_\omega \equiv \omega L$. A consumer with subjective probability $\tilde{p}_{ij} = p_i(1 + \nu_j)$ is willing to pay $A_\omega \tilde{p}_{ij}$ for coverage ω , and purchases whenever $A_\omega p_i(1 + \nu_j) \geq R$, i.e. whenever $\nu_j \geq R/(A_\omega p_i) - 1$. Demand from risk class i is therefore $D_i(R) = 1 - R/[A_\omega p_i(1 + h)]$ truncated to $[0, 1]$. This means the problem is identical in form to Proposition 5, just with A_ω replacing A and L_ω replacing L and thus the previous analytical results about the influence of risk and belief heterogeneity also apply to coinsurance for a fixed ω .

Fixing ω and substituting A_ω and L_ω for A and L throughout the proof of Proposition 5 gives a new regime threshold $\kappa^*(\omega)$. Defining $\kappa^*(\omega)$ in an analogous fashion to equation 34 above, we find that for $\kappa \leq \kappa^*(\omega)$, premiums and profits are

$$R_I^*(\omega, \kappa) = \frac{A_\omega(1 + h) + L_\omega}{2q(\kappa)} \quad (43)$$

$$\Pi_I^*(\omega, \kappa) = \frac{[A_\omega(1 + h) + L_\omega]^2}{4A_\omega(1 + h)q(\kappa)} - L_\omega p \quad (44)$$

and for $\kappa > \kappa^*(\omega)$, premiums and profits are

$$R_{II}^*(\omega, \kappa) = A_\omega(1 + h)c_M(\omega)p(1 + \kappa) \quad (45)$$

$$\Pi_{II}^*(\omega, \kappa) = \frac{p(1 + \kappa)^2(1 - c_M(\omega)) [A_\omega c_M(\omega)(1 + h) - L_\omega]}{4\kappa} \quad (46)$$

where $c_M(\omega)$ solves $B_\omega(1 - c_M(\omega)) = -c_M(\omega) \ln c_M(\omega)$ where $B_\omega \equiv [A_\omega(1 + h) + L_\omega]/[2A_\omega(1 + h)]$.

The second step in the optimal coinsurance problem is to choose the value of ω that generates the

highest profits. For each κ , this requires solving

$$\omega^*(\kappa) = \arg \max_{\omega \in [0,1]} \Pi^*(\omega, \kappa) \quad (47)$$

using whichever of $\Pi_I^*(\omega, \kappa)$ or $\Pi_{II}^*(\omega, \kappa)$ applies at that (ω, κ) pair. This problem has no closed-form solution so it was solved numerically, evaluating $\Pi^*(\omega, \kappa)$ on a fine grid of $\omega \in [0, 1]$ for each κ and taking the maximizer.

The results point to full insurance ($\omega = 1$) as the profit-maximizing policy for most of the parameter space we considered. Across a grid of loss sizes $L \in \{0.1, \dots, 0.5\}$, risk aversion parameters $\rho \in \{0.5, 1, \dots, 5\}$ and belief heterogeneity parameters $h \in \{1, \dots, 4.5\}$, full coverage is optimal for almost all combinations. Departures from full coverage only emerge as optimal when L and ρ are both near the top of this range, and only for h near 1: coverage returns to exactly 1.000 for all combinations with $h \geq 1.5$. The lowest optimal coverage rate, for $h = 1, L = 0.5, \rho = 5$, is $\omega^* \approx 0.941$.

One concern about these calculations is that, for some of the examples considered, those with high levels of both potential loss and risk aversion, the linear approximation to the utility function in equation 3 may not be as accurate. Solving the model numerically with the exact CRRA utility function rather than the linearized demand curve gives similar results. The lowest value of the coinsurance rate derived was $\omega^* \approx 0.857$, with full coverage being optimal for all values by $h \approx 2.5$. The calculations thus suggest that full insurance is therefore the optimal pooled contract, or nearly so, for realistic parameters.

7. Application: Pooling versus Risk-Based Pricing

Some risk factors, such as age, location or gender, are directly observable to insurers, so they can choose to set premiums based upon them. However, regulators can require insurers to ignore an observable factor, setting the same price for all customers. Examples include various types of community rating in health insurance and the EU's ban on gender-based pricing in auto and life insurance. Here, we assume insurers can see each group's risk type and will charge separate premiums whenever policy allows it, and ask whether mandating pooled pricing raises or lowers welfare, and how the answer depends on market structure. All derivations are in Appendices A.7 and A.8.

7.1. Two Concepts of Welfare

Under risk-based pricing, an insurer that can observe each group's risk type can charge a separate premium R_i to each group, tailored to that group's own risk. Under pooled pricing, a single premium R applies to everyone regardless of risk type. In assessing the welfare implications of these two options, you can consider two different concepts of welfare.

One approach is to follow Einav, Finkelstein and Cullen (2010) and define total welfare as the insurer's profit plus consumer surplus, with surplus calculated from willingness to pay. This is a natural choice but willingness to pay is a function of the buyer's subjectively perceived probability of loss, not the true probability. Spinnewijn (2017) shows that when willingness to pay deviates from the true value of insurance, welfare measures built on it can be systematically misleading about the impacts of policy changes.

In our case, a policy maker might instead want to know how much a buyer's insurance is genuinely worth to them: not their own willingness to pay, $p_i(1 + \nu_j)A$, but the value of the coverage evaluated at the true risk probability, p_iA . We call this the *objective* welfare criterion, and the willingness-to-pay-based measure the *subjective* criterion, and make the objective criterion our principal measure.

We use the proportional beliefs model of Section 3 to calculate objective consumer surplus and welfare under both pooled and risk-based pricing. In the interest of space, and because our earlier calculations generally pointed to this being the most realistic region, we assume here that the underlying parameter values correspond to demand from all risk groups being between zero and one.

For a single risk group facing premium R_i , the total objective consumer surplus, obtained from multiplying the surplus for each individual by total coverage, is

$$CS_i(R_i) = (Ap_i - R_i)D_i(R_i) = Ap_i - R_i - \frac{R_i}{1+h} + \frac{R_i^2}{Ap_i(1+h)} \quad (48)$$

The profit earned by the insurer from this group is

$$\Pi_i(R_i) = (R_i - Lp_i)D_i(R_i) = R_i - Lp_i + \frac{LR_i}{A(1+h)} - \frac{R_i^2}{Ap_i(1+h)} \quad (49)$$

Adding profit to consumer surplus gives objective welfare

$$W_i(R_i) = (A - L) \left[p_i - \frac{R_i}{A(1+h)} \right] \quad (50)$$

which is linear in R_i because the curvatures in consumer surplus and in profit cancel. Higher premiums for each group mean lower coverage and this reduces welfare.

Aggregating welfare, summing across the population of risk groups, does not require knowing the entire schedule of premiums R_i : the population average depends on just one summary statistic of the schedule, its mean $\mu_1 \equiv E[R_i]$.

7.2. Competition

With risk-based pricing, each group is priced at $R_i = p_i L$, giving identical demand $D^c \equiv 1 - L/[A(1+h)]$ across every group, since p_i cancels. Price equals expected cost, so profit is zero and

$$CS_C^{RB} = W_C^{RB} = p(A - L)D^c \quad (51)$$

This is independent of κ : risk-based pricing fully accommodates risk heterogeneity, so its extent is irrelevant to welfare under this benchmark.

With pooled pricing, a single price R^{PC} , derived in Proposition 6, clears the market at zero profit, so $CS = W$. Aggregating actual surplus over the population and imposing the zero-profit condition gives

$$W_C^{Pool} = (A - L) \left(p - \frac{R^{PC}}{A(1+h)} \right) \quad (52)$$

so that

$$W_C^{Pool} - W_C^{RB} = \frac{A - L}{A(1+h)} (pL - R^{PC}) \quad (53)$$

Since $R^{PC} = pL$ at $\kappa = 0$ and is strictly increasing in κ , pooling strictly reduces welfare relative to risk-based pricing whenever risk is heterogeneous, by an amount that grows with κ : adverse selection pushes the pooled price above what the population's mean risk alone would justify.

7.3. Monopoly

With risk-based pricing, the monopolist maximizes $(R_i - p_i L) D_i(R_i)$ by setting $R_i^{M, RB} = p_i \beta / 2$. Demand is

$$D^m = 1 - \beta / [2A(1 + h)] = D^c / 2 \quad (54)$$

which is the standard monopoly halving of quantity under linear demand. Total objective welfare is

$$W_M^{RB} = p(A - L) D^m = \frac{1}{2} W_C^{RB} \quad (55)$$

The optimal pooled price is Proposition 2's $R^* = \beta / [2q(\kappa)]$. Objective welfare takes the identical functional form as pooled competition, evaluated at this price:

$$W_M^{Pool} = (A - L) \left(p - \frac{\beta}{2A(1 + h)q(\kappa)} \right) \quad (56)$$

Comparing to risk-based monopoly, average welfare across all groups is

$$W_M^{Pool} - W_M^{RB} = \frac{A - L}{A(1 + h)} \left(\frac{p\beta}{2} - \frac{\beta}{2q(\kappa)} \right) \quad (57)$$

In this case, the impact of pooling on welfare is the opposite of the competitive case: since $R^* = p\beta/2$ at $\kappa = 0$ and is strictly decreasing in κ in the assumed interior region, pooling raises welfare relative to risk-based pricing under monopoly, by an amount that grows with κ . Removing the distortions associated with monopoly power provides welfare gains that strictly offset the reductions due to the pricing inefficiencies associated with different risk groups being charged the same premiums.

The increase in welfare is solely due to rising consumer surplus because the insurer's profit falls under pooling due to its reduced pricing flexibility:

$$\Pi_M^{Pool} - \Pi_M^{RB} = - \frac{(pq(\kappa) - 1) \beta^2}{4A(1 + h)q(\kappa)} \quad (58)$$

This is strictly negative for $\kappa > 0$.

These results can be explained as follows. Section 3 established that pooling and risk-based pricing produce the same total quantity of coverage; only the composition of buyers changes. At any fixed price, demand $D_i(R)$ is strictly increasing in p_i . Under risk-based pricing this is irrelevant, since every group faces its own tailored price and ends up buying at the identical rate. Under pooling, everyone faces the same price R^* , high-risk groups buy more than they would at their own fair price and low-risk groups buy less.

Total consumer surplus rises because the same quantity of coverage is now serving a riskier population, and coverage is worth more, objectively, to riskier buyers — Ap_i rises with p_i . Coverage is

also costlier to provide with pooling, since Lp_i rises with p_i too, so the insurer's profit falls. Because $A > L$, the value of serving an additional high-risk buyer always exceeds the cost of doing so, so the two effects do not offset: welfare rises.

7.4. Some Calculations of Welfare Effects

Table 1 describes how a switch from risk-based pricing to pooling affects objective consumer surplus, the insurer's profit and total objective welfare. To provide concrete numbers, we assume risk heterogeneity is determined by the uniform distribution described in Section 4. The table reports a range of values of h and κ , with the other parameters set to the values used in the rest of the paper. In each case, the parameters chosen produce the interior solution.

The figures in the table are expressed as the dollar change scaled by the same reference point, total welfare under risk-based pricing, W_M^{RB} , at that value of h . This has the advantage that the columns are additive by construction: the consumer surplus and profit columns sum exactly to the welfare column in every row. The results show that the welfare gains from pooling rise with risk heterogeneity, κ , and fall with the belief heterogeneity parameter, h , with the former effect being stronger than the latter. Gains in consumer surplus become very large for high values of κ , with gains equal to over 50% of welfare under risk-based pricing, with reductions in profits also being over 50% of this benchmark.

Table 2 shows that the losses in objective welfare implied by pooled pricing under competition are much more modest, with the maximum loss due to pooling being 1.5 percentage points for $\kappa = 0.3$ and $h = 1$. The losses also get smaller as h increases because belief heterogeneity becomes more important in driving demand relative to risk heterogeneity, thus reducing the welfare losses generated by adverse selection.

Table 1: Monopoly, objective welfare: change from switching to pooled pricing, in percentage points of risk-based welfare W_M^{RB} , at $p = 0.10$, $L = 0.20$, $\rho = 2.0$, across $h = 1.0, 2.0$, and 3.0 . Columns sum exactly: Consumer Surplus + Profit = Welfare in every row. κ is restricted to 0.30 or below so that all three panels stay within the interior regime.

κ	Consumer Surplus	Profit	Welfare
$h = 1.0$			
0.10	+6.24	−5.46	+0.78
0.20	+25.16	−22.02	+3.15
0.30	+57.41	−50.23	+7.18
$h = 2.0$			
0.10	+6.06	−5.48	+0.58
0.20	+24.45	−22.12	+2.33
0.30	+55.78	−50.47	+5.31
$h = 3.0$			
0.10	+6.52	−6.02	+0.50
0.20	+26.28	−24.26	+2.02
0.30	+59.97	−55.36	+4.61

Table 2: Competition, objective welfare: change from switching to pooled pricing, in percentage points of risk-based welfare W_C^{RB} , at $p = 0.10$, $L = 0.20$, $\rho = 2.0$. Profit is zero and consumer surplus coincides with welfare under both pricing regimes, so only welfare is shown.

κ	$h = 1.0$	$h = 2.0$	$h = 3.0$
0.10	-0.15	-0.04	-0.02
0.20	-0.62	-0.18	-0.09
0.30	-1.50	-0.43	-0.20

7.5. Subjective Welfare

The subjective welfare criteria comparisons are derived fully in Appendix A.8. Under competition, the impact of pooled pricing on welfare agrees in sign with the objective criterion result, for the same reason: risk-based pricing already prices every group at its own actuarially fair premium, so pooling can only push the price up. Under monopoly, consumers perceive the switch to pooling as being positive for consumer surplus, just as it is under the objective criterion. One difference, however, is that total subjective welfare only rises under pooling when $A(1 + h) > 3L$; when this does not hold, the fall in profit outweighs the perceived rise in consumer surplus.

8. Conclusion

It is widely understood that the adverse selection problems associated with widening the range of objective risk among customers can lead to higher prices and potential market collapse under pooled insurance.

This paper has shown that, under realistic conditions, these results may not hold. Incorporating heterogeneous beliefs about the probability of loss, we show that whether adverse selection raises premiums, leaves them unchanged, or lowers them depends on how willingness to pay is related to risk and whether insurers have market power. We have shown examples in which greater risk heterogeneity has no effect on premiums and in which it actually reduces them. We have also shown examples in which high levels of risk heterogeneity do not cause collapse, but instead lead to rising premiums and lower coverage, without ever moving the market towards no coverage. The effect of adverse selection ultimately depends on the interaction between selection, demand heterogeneity, and market structure.

Our framework explains why monopoly providers of insurance for products that have significant heterogeneity in beliefs about loss probabilities—such as appliance warranties or travel insurance—can offer pooled contracts with very high profit rates. It also provides a concrete way to assess the welfare consequences of pooled pricing policies that require insurers to charge a common premium to consumers with different risks. While pooling is inefficient when markets are competitive, we show that pooled pricing can raise both welfare and consumer surplus under monopoly because the inefficiency created by monopoly pricing exploiting consumers' mistaken beliefs about their own risk can exceed the inefficiency created by charging premiums that are not aligned with risk.

The paper's framework—in which agents differ in both economically meaningful ways but also in terms of idiosyncratic factors—may also be extended to other situations. To give two examples, workers may differ in both their objective productivity and their bilateral willingness to work for a specific firm while firms that borrow from a bank may differ in their objective probability of default while also differing in their willingness to pay for credit at various interest rates. These situations may also involve the two types of heterogeneity interacting to affect market pricing and existence in the unexpected ways demonstrated here.

References

- [1] Abito, Jose Miguel and Yuval Salant (2019). "The Effect of Product Misperception on Economic Outcomes: Evidence from the Extended Warranty Market," *Review of Economic Studies*, Volume 86, pages 2285–2318.
- [2] Akerlof, George A. (1970). "The Market for 'Lemons': Quality Uncertainty and the Market Mechanism," *Quarterly Journal of Economics*, Volume 84, pages 488–500.
- [3] Barseghyan, Levon, Francesca Molinari, Ted O'Donoghue and Joshua C. Teitelbaum (2013). "The Nature of Risk Preferences: Evidence from Insurance Choices," *American Economic Review*, Volume 103, pages 2499–2529.
- [4] Cohen, Alma and Liran Einav (2007). "Estimating Risk Preferences from Deductible Choice," *American Economic Review*, Volume 97, pages 745–788.
- [5] Corless, Robert, Gaston Gonnet, D.E.G. Hare, David Jeffrey and Donald Knuth (1996). "On the Lambert W Function," *Advances in Computational Mathematics*, Volume 5, pages 329–359.
- [6] Einav, Liran, Amy Finkelstein and Mark R. Cullen (2010). "Estimating Welfare in Insurance Markets Using Variation in Prices," *Quarterly Journal of Economics*, Volume 125, pages 877–921.
- [7] Finkelstein, Amy and Kathleen McGarry (2006). "Multiple Dimensions of Private Information: Evidence from the Long-Term Care Insurance Market," *American Economic Review*, Volume 96, pages 938–958.
- [8] Handel, Benjamin R., Jonathan T. Kolstad and Johannes Spinnewijn (2019). "Information Frictions and Adverse Selection: Policy Interventions in Health Insurance Markets," *Review of Economics and Statistics*, Volume 101, pages 326–340.
- [9] Hegarty, Tadgh and Karl Whelan (2026). "Market Structure and Prices in Online Betting Markets: Theory and Evidence," *Oxford Economic Papers*, Volume 78, pages 90–113.
- [10] Mahoney, Neale and E. Glen Weyl (2017). "Imperfect Competition in Selection Markets," *Review of Economics and Statistics*, Volume 99, pages 637–651.
- [11] Rothschild, Michael and Joseph Stiglitz (1976). "Equilibrium in Competitive Insurance Markets: An Essay on the Economics of Imperfect Information," *Quarterly Journal of Economics*, Volume 90, pages 629–649.
- [12] Spinnewijn, Johannes (2013). "Insurance and Perceptions: How to Screen Optimists and Pessimists," *Economic Journal*, Volume 123, pages 606–633.

- [13] Spinnewijn, Johannes (2017). “Heterogeneity, Demand for Insurance and Adverse Selection,” *American Economic Journal: Economic Policy*, Volume 9, pages 308–343.
- [14] Stiglitz, Joseph E. (1977). “Monopoly, Non-Linear Pricing and Imperfect Information: The Insurance Market,” *Review of Economic Studies*, Volume 44, pages 407–430.
- [15] Sydnor, Justin (2010). “(Over)insuring Modest Risks,” *American Economic Journal: Applied Economics*, Volume 2, pages 177–199.

A Omitted Proofs

A.1 Proof of Proposition 3

Zero profit under competition requires

$$\Pi(R) = R - Lp - \frac{R^2 q}{A(1+h)} + \frac{LR}{A(1+h)} = 0 \quad (\text{A.1})$$

Multiplying by $A(1+h)$,

$$qR^2 - \beta R + ALp(1+h) = 0 \quad (\text{A.2})$$

with roots

$$R = \frac{\beta \pm \sqrt{\beta^2 - 4ALp(1+h)q}}{2q} \quad (\text{A.3})$$

Since the coefficient of R^2 in the underlying profit function is negative, $\Pi(R)$ is concave: it is negative outside the two roots and strictly positive between them. Under free entry, any price in that interior range invites undercutting — a rival can always shave the price and capture the market while still earning a positive margin. The only price compatible with both entry and viability is the smaller root

$$R^{PC} = \frac{\beta - \sqrt{\beta^2 - 4ALp(1+h)q}}{2q} \quad (\text{A.4})$$

The larger root is not a stable outcome because, at that price, cutting further still earns positive profit, so competitive undercutting keeps pushing the price down past it, all the way to R^{PC} .

At $q = 1/p$,

$$R^{PC} = \frac{\beta - \sqrt{\beta^2 - 4AL(1+h)}}{2/p} = \frac{p \left[\beta - \sqrt{\beta^2 - 4AL(1+h)} \right]}{2} \quad (\text{A.5})$$

Since

$$\beta^2 - 4AL(1+h) = [A(1+h) + L]^2 - 4AL(1+h) = [A(1+h) - L]^2 \quad (\text{A.6})$$

the square root simplifies to $A(1+h) - L$, giving

$$R^{PC} = \frac{p[A(1+h) + L - A(1+h) + L]}{2} = \frac{2pL}{2} = pL \quad (\text{A.7})$$

For monotonicity of premiums as a function of q , write the quadratic as an identity in $R(q)$,

$$q[R(q)]^2 - \beta R(q) + ALp(1+h) = 0 \quad (\text{A.8})$$

which holds for every q . Differentiating with respect to q and setting the result to zero,

$$R^2 + (2qR - \beta) \frac{dR}{dq} = 0 \Rightarrow \frac{dR}{dq} = \frac{R^2}{\beta - 2qR} \quad (\text{A.9})$$

By Vieta's formula, the two roots of $qR^2 - \beta R + ALp(1+h) = 0$ sum to β/q . Since R^{PC} is the smaller of the two, it is below their average, $R^{PC} < \beta/(2q)$, i.e. $\beta - 2qR^{PC} > 0$, and hence $dR^{PC}/dq > 0$.

For the final part of the proof, the discriminant in A.4 is non-negative provided $q \leq \beta^2/[4ALp(1+h)] = q_0$, as required. ■

A.2 Proof of Proposition 4

Aggregate demand is $D(R) = 1 - Rq/[A(1+h)]$, so inverse demand is $R(Q) = A(1+h)(1-Q)/q$. Total claims cost, as a function of R , is

$$T(R) = E[p_i D(p_i, R)] = p - \frac{R}{A(1+h)} \quad (\text{A.10})$$

With cost shared in proportion to output share, firm f 's profit is

$$\pi_f = R(Q)q_f - \frac{q_f}{Q} L T(R(Q)) \quad (\text{A.11})$$

where

$$Q = q_f + Q_{-f} \quad (\text{A.12})$$

Differentiating with respect to q_f and imposing symmetry ($Q_{-f} = (n-1)q_f$, $Q = nq_f$) gives a quadratic in aggregate Q :

$$A(1+h)(n+1)Q^2 - n[A(1+h) - L]Q + L(n-1)(pq - 1) = 0 \quad (\text{A.13})$$

Taking the root with the plus sign gives Q_n^* as stated. At $q = 1/p$, the discriminant simplifies to a perfect square,

$$n^2[A(1+h) - L]^2 - 4A(1+h)(n+1)L(n-1)(0) = n^2[A(1+h) - L]^2 \quad (\text{A.14})$$

given the homogeneous-risk solution

$$Q_0^* = n[A(1+h) - L]/\{A[h(n+1) + n + 1]\} \quad (\text{A.15})$$

At $n = 1$, the discriminant simplifies to $[A(1+h) - L]^2$ regardless of q (the factor $n - 1 = 0$ eliminates the q -dependence entirely), and Q_n^* reduces to

$$Q_1^* = [A(1+h) - L]/[2A(1+h)] \quad (\text{A.16})$$

the monopoly quantity in Proposition 2.

Existence. A real root of the quadratic is necessary but not sufficient for a genuine equilibrium: firms will not operate at a price the quadratic identifies if the resulting profit is negative, since producing zero is always available. Industry profit at any price R is

$$\Pi(R) = R[1 - Rq/(A(1+h))] - L[p - R/(A(1+h))] \quad (\text{A.17})$$

the same function whose maximum over R is $\Pi_I^*(q)$, zero exactly at $q = q_0$ and negative for $q > q_0$. Since $\Pi_I^*(q)$ is the maximum of $\Pi(R)$ over every price, no price — including whatever the Cournot quadratic identifies — can yield positive industry profit once $q > q_0$. The quadratic's own discriminant remains real up to a strictly larger threshold,

$$q_{max}(n) = 1/p + [A(1+h) - L]^2 n^2 / [4ALp(1+h)(n^2 - 1)] > q_0 \quad (\text{A.18})$$

for every finite n , so for $q \in (q_0, q_{max}(n))$ the quadratic still returns a mathematical root, but one at which industry — and by symmetry, each firm's own — profit is strictly negative: not a viable equilibrium.

The correct existence threshold is q_0 itself, for every n . At $q = q_0$, substituting $q_0 = \beta^2/[4ALp(1+h)]$ into the quadratic shows $Q = Q_I^* = [A(1+h) - L]/[2A(1+h)]$ solves it exactly, independent of n : the n -firm Cournot equilibrium coincides with the monopoly quantity itself at this boundary, not merely with the same profit level, which is why profit is exactly zero there for every n . At $q = 1/p$, industry profit simplifies to $np[A(1+h) - L]^2/[A(1+h)(n+1)^2]$, which is positive since $A(1+h) > L$ always.

We can now show industry profit is non-negative throughout $q \in (1/p, q_0]$, for every n . Define $z \equiv pq$, so q_0 corresponds to

$$z_0 \equiv pq_0 = 1 + [A(1+h) - L]^2/[4A(1+h)L] \quad (\text{A.19})$$

and $q \leq q_0$ is equivalent to

$$[A(1+h) - L]^2 - 4A(1+h)L(z-1) \geq 0 \quad (\text{A.20})$$

Industry profit at any quantity Q is

$$\Pi(Q) = R(Q)Q - LT(Q) \quad (\text{A.21})$$

with

$$R(Q) = A(1+h)(1-Q)/q \quad (\text{A.22})$$

and

$$T(Q) = p - (1-Q)/q \quad (\text{A.23})$$

which simplifies to

$$q\Pi(Q) = (1-Q)[A(1+h)Q + L] - Lz \quad (\text{A.24})$$

For $n > 1$, solving the equilibrium quadratic for z and substituting eliminates z entirely:

$$q\Pi(Q_n^*) = \frac{Q_n^* \{2A(1+h)Q_n^* - [A(1+h) - L]\}}{n-1} \quad (\text{A.25})$$

Writing

$$S \equiv \sqrt{n^2[A(1+h) - L]^2 - 4A(1+h)(n+1)(n-1)L(z-1)} \quad (\text{A.26})$$

so that

$$Q_n^* = \{n[A(1+h) - L] + S\} / [2A(1+h)(n+1)] \quad (\text{A.27})$$

substitution gives

$$2A(1+h)Q_n^* - [A(1+h) - L] = \frac{S - [A(1+h) - L]}{n+1} \quad (\text{A.28})$$

and

$$q\Pi(Q_n^*) = \frac{Q_n^* \{S - [A(1+h) - L]\}}{n^2 - 1} \quad (\text{A.29})$$

Finally, we need to establish the sign of $S - [A(1+h) - L]$. Since

$$S^2 - [A(1+h) - L]^2 = (n^2 - 1) \{[A(1+h) - L]^2 - 4A(1+h)L(z-1)\} \quad (\text{A.30})$$

condition (A.20) gives $S^2 \geq [A(1+h) - L]^2$, and since both $S \geq 0$ and $A(1+h) - L \geq 0$, $S \geq A(1+h) - L$. With $Q_n^* > 0$ and $n^2 - 1 > 0$, every factor in $q\Pi(Q_n^*)$ is non-negative, so $\Pi(Q_n^*) \geq 0$ for every $n > 1$ and every $q \leq q_0$.

At $n = 1$ the equilibrium quadratic reduces to $2A(1+h)Q^2 - [A(1+h) - L]Q = 0$, which solves to give the monopoly quantity derived earlier. Substituting into $q\Pi(Q)$ gives $q\Pi(Q_1^*) = \beta^2/[4A(1 +$

$h)] - Lz$, non-negative exactly when $z \leq \beta^2/[4A(1+h)L] = z_0$ — the same condition as for $n > 1$. Industry profit is therefore non-negative throughout $q \in (1/p, q_0]$ for every $n \geq 1$.

Convergence. For the limiting premium as n increases, define $\beta' \equiv A(1+h) - L$ and factor the discriminant as

$$n^2\beta'^2 - 4A(1+h)L(n^2-1)(pq-1) = n^2[\beta'^2 - 4A(1+h)L(pq-1)] + 4A(1+h)L(pq-1) \quad (\text{A.31})$$

As $n \rightarrow \infty$, the bracketed term dominates and the additive remainder becomes negligible relative to it, so

$$\sqrt{n^2[A(1+h) - L]^2 - 4A(1+h)L(n+1)L(n-1)(pq-1)} \sim n\sqrt{\beta'^2 - 4A(1+h)L(pq-1)} \quad (\text{A.32})$$

Substituting this into Q_n^* and using $n/(n+1) \rightarrow 1$,

$$Q_n^* \rightarrow \frac{\beta' + \sqrt{\beta'^2 - 4A(1+h)L(pq-1)}}{2A(1+h)} \quad (\text{A.33})$$

Expanding

$$\beta'^2 - 4A(1+h)L(pq-1) = [A(1+h) - L]^2 - 4A(1+h)L(pq-1) = \beta^2 - 4ALp(1+h)q \quad (\text{A.34})$$

this limit is exactly

$$Q_n^* \rightarrow \frac{A(1+h) - L + \sqrt{\beta^2 - 4ALp(1+h)q}}{2A(1+h)} \quad (\text{A.35})$$

which is the quantity implied by Proposition 3's competitive premium, $Q^{PC} = 1 - R^{PC}q/[A(1+h)]$, verified by direct substitution of $R^{PC} = [\beta - \sqrt{\beta^2 - 4ALp(1+h)q}]/(2q)$ into that expression and simplifying. Since $R_n^* = A(1+h)(1 - Q_n^*)/q$ is a continuous function of Q_n^* , $Q_n^* \rightarrow Q^{PC}$ implies $R_n^* \rightarrow R^{PC}$ directly. ■

A.3 Proof of Proposition 5

The monopolist's profit is

$$\Pi(x) = \frac{1}{2\kappa} \int_{-\kappa}^{\kappa} (Ax - Lp_i) D(\delta_i, x) d\delta_i \quad (\text{A.36})$$

Interior piece. For any premium x such that $x \leq p(1-\kappa)(1+h)$, every risk type has positive demand strictly below 1, so the bounds in $D(\delta_i, x)$ do not bind and $D(\delta_i, x) = 1 - x/[p_i(1+h)]$ over the entire range of integration. This case is Proposition 2 above, specialized to $q = q(\kappa)$, giving $x_I^* = [A(1+h) + L]/(2Aq(\kappa))$ and $R_I^* = Ax_I^* = [A(1+h) + L]/(2q(\kappa))$.

Truncation piece. For $x > p(1-\kappa)(1+h)$, the lower bound binds for the lowest-risk types: $D(\delta_i, x) =$

0 for $\delta_i < \bar{\delta}(x) \equiv x/[p(1+h)] - 1$, and the integral runs only over $[\bar{\delta}(x), \kappa]$:

$$\Pi(x) = \frac{1}{2\kappa(1+h)} \int_{\bar{\delta}(x)}^{\kappa} (Ax - Lp_i) \left(1 + h - \frac{x}{p_i}\right) d\delta_i \quad (\text{A.37})$$

Differentiating with respect to x , both the lower limit $\bar{\delta}(x)$ and the integrand depend on x . The full Leibniz derivative is:

$$\begin{aligned} \frac{d\Pi}{dx} &= \frac{1}{2\kappa(1+h)} \left[\int_{\bar{\delta}}^{\kappa} \frac{\partial}{\partial x} \left[(Ax - Lp_i) \left(1 + h - \frac{x}{p_i}\right) \right] d\delta_i \right. \\ &\quad \left. - (Ax - Lp(\bar{\delta})) \left(1 + h - \frac{x}{p(1+\bar{\delta})}\right) \frac{d\bar{\delta}}{dx} \right] \end{aligned} \quad (\text{A.38})$$

The boundary term vanishes: by the definition of $\bar{\delta}$, $x/[p(1+\bar{\delta})] = 1+h$, so $1+h-x/p(1+\bar{\delta}) = 0$. The derivative inside the integral is:

$$\begin{aligned} \frac{\partial}{\partial x} \left[(Ax - Lp_i) \left(1 + h - \frac{x}{p_i}\right) \right] &= A \left(1 + h - \frac{x}{p_i}\right) + (Ax - Lp_i) \left(-\frac{1}{p_i}\right) \\ &= A(1+h) - \frac{Ax}{p_i} - \frac{Ax}{p_i} + L \\ &= A(1+h) + L - \frac{2Ax}{p_i} \end{aligned} \quad (\text{A.39})$$

Setting $d\Pi/dx = 0$ and multiplying through by the constant $(1+h)$:

$$\int_{\bar{\delta}(x)}^{\kappa} \left[A(1+h) + L - \frac{2Ax}{p(1+\delta_i)} \right] d\delta_i = 0 \quad (\text{A.40})$$

An integral calculation required here is

$$\int_{\bar{\delta}}^{\kappa} \frac{d\delta_i}{1+\delta_i} = \ln\left(\frac{1+\kappa}{1+\bar{\delta}}\right) \quad (\text{A.41})$$

The first-order condition therefore becomes:

$$[A(1+h) + L](\kappa - \bar{\delta}) = \frac{2Ax}{p} \ln\left(\frac{1+\kappa}{1+\bar{\delta}}\right) \quad (\text{A.42})$$

Both terms in this condition involve $\bar{\delta}$ only relative to κ , through $\kappa - \bar{\delta}$ and $\ln[(1+\kappa)/(1+\bar{\delta})]$, suggesting the marginal buyer's boundary is best expressed as a fixed fraction of the population's ceiling $1+\kappa$. Defining c_M as this ratio,

$$c_M \equiv \frac{1+\bar{\delta}}{1+\kappa}, \quad c_M \in (0, 1) \quad (\text{A.43})$$

gives

$$\kappa - \bar{\delta} = (1 - c_M)(1 + \kappa) \quad (\text{A.44})$$

and

$$\ln[(1 + \kappa)/(1 + \bar{\delta})] = -\ln c_M \quad (\text{A.45})$$

The same substitution, via the definition $\bar{\delta}(x) \equiv x/[p(1 + h)] - 1$, pins down the candidate premium itself:

$$x = p(1 + h)c_M(1 + \kappa) \quad (\text{A.46})$$

so

$$2Ax/p = 2A(1 + h)c_M(1 + \kappa) \quad (\text{A.47})$$

Substituting all three pieces into the first-order condition gives

$$[A(1 + h) + L](1 - c_M)(1 + \kappa) = 2A(1 + h)c_M(1 + \kappa)(-\ln c_M) \quad (\text{A.48})$$

Canceling $(1 + \kappa) > 0$ and substituting $B \equiv [A(1 + h) + L]/[2A(1 + h)]$ (so $A(1 + h) + L = 2A(1 + h)B$):

$$2A(1 + h)B(1 - c_M) = 2A(1 + h)c_M(-\ln c_M) \quad (\text{A.49})$$

Dividing through by $2A(1 + h)$ gives $B(1 - c_M) = -c_M \ln c_M$.

Existence and uniqueness of c_M . Let

$$\phi(c) \equiv B(1 - c) + c \ln c \quad (\text{A.50})$$

so that c_M which solves $B(1 - c_M) = -c_M \ln c_M$ is equivalent to $\phi(c_M) = 0$.

Direct substitution gives that $\phi(1) = 0$: $c = 1$ is always a root, corresponding to the trivial case in which no truncation occurs. As $c \rightarrow 0^+$, $\phi(c) \rightarrow B > 0$. And $\phi'(1) = 1 - B$, which is strictly positive whenever $A(1 + h) > L$ since this makes $B = \frac{1}{2} + L/[2A(1 + h)] < 1$; so ϕ is increasing at $c = 1$, meaning $\phi(c) < 0$ immediately to the left of 1. Combined with $\phi(0^+) = B > 0$, the intermediate value theorem gives a root strictly inside $(0, 1)$. Since $\phi''(c) = 1/c > 0$ for all $c > 0$, ϕ is strictly convex, and a strictly convex function can cross zero at most twice; having already done so once at $c = 1$, there is exactly one further root, which is c_M .

This root has a closed form, $c_M = -B/W_{-1}(-Be^{-B})$, where W_{-1} is the lower real branch of the Lambert W function (Corless et al., 1996); the other real branch, W_0 , returns the trivial root $c = 1$.

Solution with truncation. Since $p_i = p(1 + \delta_i)$ and $1 + \bar{\delta} = c_M(1 + \kappa)$, the marginal purchasing type's risk level is $\bar{p} = c_M p(1 + \kappa)$, and the solution on this piece is

$$x_{II}^* = (1 + h)c_M p(1 + \kappa) \quad (\text{A.51})$$

$$R_{II}^* = Ax_{II}^* = A(1 + h)c_M p(1 + \kappa) \quad (\text{A.52})$$

Continuity. At $\kappa = \kappa^*$, the interior solution x_I^* satisfies $x_I^* = p(1 - \kappa)(1 + h)$ by definition of the threshold: this is exactly the boundary at which the lowest-risk type's demand reaches zero. At this same point, the lower-truncation regime's marginal purchasing type $\bar{p}(x) = x/(1 + h)$ equals $p(1 - \kappa)$, so the truncation regime's domain of integration $[\bar{\delta}(x), \kappa]$ converges to the full support $[-\kappa, \kappa]$ used in the interior regime. Since both pieces of $\Pi(x)$ are defined by the same integrand $(Ax - Lp_i)(1 + h - x/p_i)$, merely over domains that coincide in the limit, the two profit functions agree at this shared boundary point.

Existence and uniqueness of κ^* . By definition, κ^* is the value of κ at which the interior threshold equals the price that excludes the lowest-risk type:

$$x_I^*(\kappa^*) = p(1 - \kappa^*)(1 + h) \quad (\text{A.53})$$

The right-hand side, $p(1 - \kappa)(1 + h)$, is the exact price at which the cheapest-to-insure consumer in the population — the lowest-risk type, with $p_i = p(1 - \kappa)$ — would stop buying altogether: demand $D(\delta_i, x) = 1 - x/[p_i(1 + h)]$ reaches zero exactly when $x = p_i(1 + h)$. As risk heterogeneity κ rises, the interior premium falls (Jensen's inequality, as in Proposition 2 above), but the threshold $p(1 - \kappa)(1 + h)$ falls faster, so the two must eventually cross. κ^* is exactly that crossing point.

Substituting $x_I^*(\kappa) = [A(1 + h) + L]/(2Aq(\kappa))$ into the defining condition, then cancelling p and clearing denominators, gives

$$\kappa^* [A(1 + h) + L] = A(1 - \kappa^*)(1 + h) \ln \left(\frac{1 + \kappa^*}{1 - \kappa^*} \right) \quad (\text{A.54})$$

Define

$$f(\kappa) \equiv A(1 - \kappa)(1 + h) \ln \left(\frac{1 + \kappa}{1 - \kappa} \right) - \kappa [A(1 + h) + L] \quad (\text{A.55})$$

so that κ^* is a root of f . $f(0) = 0$ trivially. Near $\kappa = 0$, using $\ln[(1 + \kappa)/(1 - \kappa)] \approx 2\kappa$ for small κ ,

$$f(\kappa) \approx \kappa [2A(1 + h) - A(1 + h) - L] = \kappa [A(1 + h) - L] \quad (\text{A.56})$$

which is positive whenever $A(1 + h) > L$ — the same condition that makes full insurance worth offering at all. So f is positive just above $\kappa = 0$. As $\kappa \rightarrow 1^-$, the term $(1 - \kappa) \ln[(1 + \kappa)/(1 - \kappa)] \rightarrow 0$ while $-\kappa [A(1 + h) + L] \rightarrow -[A(1 + h) + L] < 0$, so $f(\kappa) \rightarrow -[A(1 + h) + L] < 0$. By continuity, f has at least one root in $(0, 1)$.

Differentiating f twice,

$$f''(\kappa) = \frac{4A(1 + h)}{(\kappa - 1)(\kappa + 1)^2} \quad (\text{A.57})$$

For $\kappa \in (0, 1)$, $\kappa - 1 < 0$ while $(\kappa + 1)^2 > 0$ and $A(1 + h) > 0$, so $f''(\kappa) < 0$ throughout: f is strictly concave on $(0, 1)$, so f' is strictly decreasing there and can equal zero at most once. This means the

solution described above is unique. ■

A.4 Implications of Proposition 5

Corollary 1. *For $\kappa \leq \kappa^*$, $R^*(\kappa)$ is strictly decreasing in κ .*

Proof. $R_I^*(\kappa) = [A(1+h) + L]/(2q(\kappa))$ is decreasing in q , so it suffices that $q = E[1/p_i]$ is strictly increasing in κ . Since $1/p_i$ is strictly convex in p_i , a mean-preserving spread of p_i — an increase in κ holding p fixed — strictly raises $E[1/p_i]$ by Jensen's inequality. ■

Corollary 2. *For $\kappa > \kappa^*$, $R^*(\kappa)$ is linear and strictly increasing in κ , coinciding with the interior branch at $\kappa = \kappa^*$.*

Proof. $B \equiv [A(1+h) + L]/[2A(1+h)]$ does not depend on κ , so c_M — the solution of $B(1 - c_M) = -c_M \ln c_M$ — does not depend on κ either. Hence $R_{II}^*(\kappa) = A(1+h)c_M p(1+\kappa)$ is linear and strictly increasing in κ . Combined with Corollary 1 and continuity at κ^* , $R^*(\kappa)$ falls on $(0, \kappa^*)$ and rises on $(\kappa^*, 1)$, with its minimum at κ^* . ■

Corollary 3. *The level of $R^*(\kappa)$ is continuous at κ^* (established in the Proposition), but its slope is not: the left- and right-derivatives differ. Maximized profit $\Pi^*(\kappa)$, by contrast, is continuous in both level and slope at κ^* — an envelope-theorem result: since x_I^* and x_{II}^* coincide at the boundary and the two regimes share the same profit integrand there, the direct effect of κ on profit at the optimum is identical on both sides, even though the optimal premium's own sensitivity to κ is not.*

Proof. On the interior side, $R_I^* = [A(1+h) + L]/(2q)$ depends on κ entirely through $q(\kappa)$. On the truncation side, $R_{II}^* = A(1+h)c_M p(1+\kappa)$ depends on κ entirely through the population's highest risk level, $p(1+\kappa)$, since c_M does not vary with κ (Corollary 2), and $dR_{II}^*/d\kappa = A(1+h)c_M p$ is a constant. These are two structurally different functions of κ , each continuous and differentiable on its own domain, that happen to take the same value at κ^* but have no reason to share a derivative there.

Profit's slope, by contrast, is continuous. Since $\Pi^*(\kappa)$ is profit maximized over x , the envelope theorem implies $d\Pi^*/d\kappa$ depends only on the direct effect of κ on the profit integrand at the optimal x^* , not on how x^* responds to κ and the direct effect is identical on both sides of κ^* . ■

Corollary 4. *Proposition 2 already showed that pooled monopoly quantity does not depend on q , hence not on κ , throughout the interior regime. Once truncation excludes the lowest-risk types, this no longer holds:*

$$Q_{II}^*(\kappa) = \frac{(1+\kappa)[(1-c_M) + c_M \ln c_M]}{2\kappa} \quad (\text{A.58})$$

and quantity falls in κ .

Proof. Substituting $x_{II}^*(\kappa)$ into $D(\delta_i, x)$ and averaging only over the purchasing types, $\delta_i \in [\bar{\delta}(\kappa), \kappa]$, gives $Q_{II}^*(\kappa) = (1+\kappa)[(1-c_M) + c_M \ln c_M]/(2\kappa)$, which falls in κ . This agrees with the interior value

from Proposition 2 at κ^* , by the same shared-integrand argument used to establish continuity of R^* and Π^* above. ■

Corollary 5. *The profit rate on premiums actually accepted, $\Pi_{II}^*(\kappa)/[R_{II}^*(\kappa)Q_{II}^*(\kappa)]$, does not depend on κ for $\kappa > \kappa^*$:*

$$\frac{\Pi_{II}^*(\kappa)}{R_{II}^*(\kappa)Q_{II}^*(\kappa)} = \frac{(1 - c_M) [Ac_M(1 + h) - L]}{2Ac_M(1 + h) [(1 - c_M) + c_M \ln c_M]} \quad (\text{A.59})$$

Proof. This follows from how the truncated regime scales with κ . The marginal purchasing type's risk level is pinned to $c_M p(1 + \kappa)$ — a fixed fraction c_M of the population's highest possible risk level $p(1 + \kappa)$, and that fraction does not itself move with κ (Corollary 2). As κ grows, the truncated regime therefore does not change shape; the whole problem rescales with the population's upper risk bound. The algebra confirms this directly. Multiplying premium and quantity,

$$\begin{aligned} R_{II}^*(\kappa)Q_{II}^*(\kappa) &= A(1 + h)c_M p(1 + \kappa) \frac{(1 + \kappa) [(1 - c_M) + c_M \ln c_M]}{2\kappa} \\ &= \frac{Ac_M p(1 + h)(1 + \kappa)^2 [(1 - c_M) + c_M \ln c_M]}{2\kappa} \end{aligned} \quad (\text{A.60})$$

and, from the definition of profit on this piece,

$$\Pi_{II}^*(\kappa) = \frac{p(1 - c_M) [Ac_M(1 + h) - L] (1 + \kappa)^2}{4\kappa} \quad (\text{A.61})$$

Both expressions carry the identical factor $p(1 + \kappa)^2/\kappa$, which cancels exactly on division:

$$\frac{\Pi_{II}^*(\kappa)}{R_{II}^*(\kappa)Q_{II}^*(\kappa)} = \frac{(1 - c_M) [Ac_M(1 + h) - L] / 4}{Ac_M(1 + h) [(1 - c_M) + c_M \ln c_M] / 2} = \frac{(1 - c_M) [Ac_M(1 + h) - L]}{2Ac_M(1 + h) [(1 - c_M) + c_M \ln c_M]} \quad (\text{A.62})$$

which is the stated formula, with no κ remaining anywhere in it. ■

A.5 Proof of Proposition 6

$\kappa^{trunc,C}$ has no closed form. Solving it numerically for $p = 0.10$, $L = 0.05$, $\rho = 2.0$ gives $\kappa^{trunc,C} \approx 0.474$ ($h = 1$), rising to ≈ 0.810 ($h = 4.5$). For $\kappa > \kappa^{trunc,C}$, the relevant profit function is exactly $\Pi(x)$ from the truncation piece of Proposition 5's proof, here set to zero rather than maximized:

$$\int_{\bar{\delta}(x)}^{\kappa} (Ax - Lp_i) \left(1 + h - \frac{x}{p_i}\right) d\delta_i = 0 \quad (\text{A.63})$$

Substituting the candidate solution $x = (1 + h)cp(1 + \kappa)$ and simplifying, using $L = (2B - 1)A(1 + h)$, gives an equation independent of κ :

$$2c^2 \ln c - (2B + 1)c^2 + 4Bc - 2B + 1 = 0 \quad (\text{A.64})$$

This does not have a closed-form solution but $c = 1$ is always a root: this is the trivial solution at which $\bar{\delta} = \kappa$, so no one buys and both revenue and cost are zero. The genuine solution is a second root $c_C \in (0, 1)$, giving $R^{PC}(\kappa) = A(1+h)c_C p(1+\kappa)$ as stated.

Existence. The maximum profit achievable in the truncated regime is $\Pi_{II}^*(\kappa)$ from Proposition 5 — the monopolist's truncated-regime profit, using monopoly's constant c_M (from $B(1 - c_M) = -c_M \ln c_M$), since it is the same profit function. Whenever $Ac_M(1+h) > L$, $\Pi_{II}^*(\kappa) > 0$ for every $\kappa \in (0, 1)$ (Proposition 5). Since the truncated profit function is continuous, equals zero at $c = 1$, and rises to a strictly positive maximum before returning to zero, it crosses zero exactly once on the ascending side, at $R^{PC}(\kappa)$, for every κ .

Fraction Insured. The marginal participating type is $\bar{\delta}(\kappa) = c_C(1+\kappa) - 1$, giving

$$\text{Share}(\kappa) = (\kappa - \bar{\delta}(\kappa))/(2\kappa) = (1+\kappa)(1-c_C)/(2\kappa) \quad (\text{A.65})$$

Evaluating this at $\kappa = 1$ gives $1 - c_C$, which is strictly positive since $c_C < 1$. ■

A.6 Proof of Proposition 7

As with the truncated regime of Propositions 5 and 6, the marginal purchasing type is pinned down by $1 + \bar{\delta} = c(1+\kappa)$ for some $c \in (0, 1)$, giving $x(c) = (1+h)cp(1+\kappa)$ and $R(c) = Ax(c)$. Integration gives

$$Q(c) = \frac{(1+\kappa)[(1-c) + c \ln c]}{2\kappa} \quad (\text{A.66})$$

$$\text{TotalCost}(c) = \frac{Lp(1+\kappa)^2(1-c)^2}{4\kappa} \quad (\text{A.67})$$

the first matching $Q_{II}^*(\kappa)$ from Proposition 5 exactly, as it must, since aggregate demand does not depend on market structure. Firm f 's profit is

$$\pi_f = q_f R(c) - (q_f/Q(c)) \text{TotalCost}(c) \quad (\text{A.68})$$

with c determined implicitly by

$$Q(c) = q_f + Q_{-f} \quad (\text{A.69})$$

Differentiating with respect to q_f via the implicit function theorem and imposing symmetry ($Q(c) = nq_f$) gives a transcendental first-order condition in c . The factor $(1-c) + c \ln c$ divides out, corresponding to the trivial equilibrium in which no one buys. The remaining factor, after substituting $L = (2B-1)A(1+h)$ and dividing by $A(1+h)$, is

$$n \ln(c) [2c^2 \ln c - (2B+1)c^2 + 4Bc - 2B + 1] + \text{Rem}(c) = 0 \quad (\text{A.70})$$

where

$$\text{Rem}(c) \equiv 4B(c-1)^2 + 2B \ln(c)(1-c^2) + 2c^2 \ln^2 c - \ln(c)(3c-1)(c-1) \quad (\text{A.71})$$

This does not have a closed-form solution but at $n = 1$ it reduces to $B(1-c) = -c \ln c$, which is Proposition 5's monopoly condition.

Dividing by n and taking $n \rightarrow \infty$, the term $\text{Rem}(c)/n \rightarrow 0$, leaving

$$\ln(c) [2c^2 \ln c - (2B+1)c^2 + 4Bc - 2B + 1] = 0 \quad (\text{A.72})$$

since $\ln(c) \neq 0$ for $c \neq 1$, this requires $2c^2 \ln c - (2B+1)c^2 + 4Bc - 2B + 1 = 0$, exactly Proposition 6's competitive equation. Both limits were additionally verified numerically: at $h = 1, p = 0.10, L = 0.05$ and $\rho = 2.0$, $c_n \rightarrow 0.5594$ (monopoly's own value) as $n \rightarrow 1$ and $c_n \rightarrow 0.3566$ (matching c_C) as $n \rightarrow 100,000$. ■

A.7 Objective Welfare Measures

A consumer (i, j) has true risk $p_i = p(1+\delta_i)$ and belief $p_i(1+\nu_j)$, $\nu_j \sim U[-1, h]$, and buys full coverage at premium R_i whenever $\nu_j > R_i/(Ap_i) - 1$, where R_i may differ across risk groups, as under risk-based pricing, or be the same for every group, as under pooled pricing. The objective welfare metric evaluates a buyer's surplus at true risk: if they buy, their contribution to social surplus is $Ap_i - R_i$, and 0 otherwise, regardless of the belief that led them to purchase. We will define $\theta \equiv A(1+h) - L$.

A.7.1 Demand and Aggregate Welfare Criteria

A buyer in group i purchases whenever $Ap_i(1+\nu_j) \geq R_i$, i.e. whenever $\nu_j \geq R_i/(Ap_i) - 1$. The fraction of group i that buys is therefore

$$D_i(R_i) = \int_{R_i/(Ap_i)-1}^h \frac{d\nu}{1+h} = 1 - \frac{R_i}{Ap_i(1+h)} \quad (\text{A.73})$$

Averaging the per-group welfare formulas of Section 5.1 over the population of risk groups with $\delta_i \sim U[-\kappa, \kappa]$ requires the following two moments

$$\mu_1 \equiv E[R_i], \quad \mu_2 \equiv E\left[\frac{R_i^2}{p_i}\right] \quad (\text{A.74})$$

both taken as averages across risk groups, not across individual buyers. In terms of μ_1 and μ_2 alone, aggregate consumer surplus, profit, and welfare are

$$CS(\mu_1, \mu_2) = Ap - \mu_1 - \frac{\mu_1}{1+h} + \frac{\mu_2}{A(1+h)} \quad (\text{A.75})$$

$$\Pi(\mu_1, \mu_2) = \mu_1 - \frac{\mu_2}{A(1+h)} - Lp + \frac{L\mu_1}{A(1+h)} \quad (\text{A.76})$$

$$W(\mu_1, \mu_2) = CS(\mu_1, \mu_2) + \Pi(\mu_1, \mu_2) = (A - L) \left[p - \frac{\mu_1}{A(1+h)} \right] \quad (\text{A.77})$$

Objective welfare, (A.77), depends only on μ_1 : aggregate objective welfare is linear in the mean premium alone, with μ_2 — and hence the dispersion or correlation structure of R_i across groups — mattering only for how this total is split between consumer surplus and profit, not for its size. Setting $\Pi(\mu_1, \mu_2) = 0$ and $\mu_1 = R^{PC}$ recovers W_C^{Pool} under pooled competition; setting $\mu_1 = p\beta/2$ recovers W_M^{RB} under risk-based monopoly; and so on for the other cases reported in the text.

A.7.2 Profit Comparison Under Monopoly

Under pooled monopoly, the premium is $R^* = \beta/[2q(\kappa)]$, while under risk-based pricing the monopolist charges $R_i = p_i\beta/2$. The corresponding values of μ_2 are

$$\mu_2^{Pool} = \frac{\beta^2}{4q(\kappa)} \quad \mu_2^{RB} = \frac{p\beta^2}{4} \quad (\text{A.78})$$

so that profits under pooling minus profits under risk-based pricing, using $\Pi(\mu_1, \mu_2)$ above, are

$$\Pi_M^{Pool} - \Pi_M^{RB} = -\frac{(pq(\kappa) - 1)\beta^2}{4A(1+h)q(\kappa)} \quad (\text{A.79})$$

which is strictly negative for $\kappa > 0$ in the interior region. The monopolist loses profits due to lower pricing flexibility. This shows that the gain in welfare reported in the text is all due to the increase in consumer surplus offsetting falling profits.

A.8 Subjective Welfare Measures

An alternative approach is to assess subjective welfare, evaluating a buyer's surplus using their own belief: if they buy, their contribution to the total subjective surplus is $Ap_i(1 + \nu_j) - R_i$.

A.8.1 Welfare Criteria by Risk Group

Integrating perceived surplus $Ap_i(1 + \nu) - R_i$ over buyers,

$$\begin{aligned}
 CS_i^{perceived}(R_i) &= \int_{R_i/(Ap_i)-1}^h \left(\frac{Ap_i(1 + \nu) - R_i}{1 + h} \right) d\nu \\
 &= \frac{1}{1 + h} \left[Ap_i\nu + \frac{Ap_i\nu^2}{2} - R_i\nu \right]_{R_i/(Ap_i)-1}^h \\
 &= \frac{1}{1 + h} \left[\left(Ap_i h + \frac{Ap_i h^2}{2} - R_i h \right) - \left(-\frac{Ap_i}{2} - \frac{R_i^2}{2Ap_i} + R_i \right) \right] \\
 &= \frac{Ap_i(1 + h)}{2} - R_i + \frac{R_i^2}{2Ap_i(1 + h)}
 \end{aligned} \tag{A.80}$$

Adding profit, $\Pi_i(R_i)$, from Section 5.1, subjective welfare is

$$W_i^{perceived}(R_i) = CS_i^{perceived}(R_i) + \Pi_i(R_i) = p_i \left[\frac{A(1 + h)}{2} - L \right] + \frac{LR_i}{A(1 + h)} - \frac{R_i^2}{2Ap_i(1 + h)} \tag{A.81}$$

which is strictly concave in R_i .

A.8.2 Aggregate Welfare Criteria

Averaging over the population of risk groups requires the same two moments, $\mu_1 \equiv E[R_i]$ and $\mu_2 \equiv E[R_i^2/p_i]$, defined in Appendix A.7.1. In terms of μ_1 and μ_2 , aggregate subjective consumer surplus and welfare are

$$CS^{perceived}(\mu_1, \mu_2) = \frac{Ap(1 + h)}{2} - \mu_1 + \frac{\mu_2}{2A(1 + h)} \tag{A.82}$$

$$W^{perceived}(\mu_1, \mu_2) = CS^{perceived}(\mu_1, \mu_2) + \Pi(\mu_1, \mu_2) \tag{A.83}$$

using $\Pi(\mu_1, \mu_2)$ from Appendix A.7.1. Unlike the objective case, $W^{perceived}$ depends on both μ_1 and μ_2 : because the subjective welfare function is concave rather than linear in the premium, the dispersion of R_i across risk groups matters for the total, not only for how it is split between consumers and the insurer.

A.8.3 Pooled and Risk-Based Pricing

Under pooled pricing, $\mu_1 = R$ and $\mu_2 = R^2 q(\kappa)$; substituting $R = R^{PC}$ (Proposition 6) or $R = R^* = \beta/[2q(\kappa)]$ (Proposition 2) gives the pooled competitive and monopoly results used below. Under risk-based competition, $R_i = p_i L$ gives $\mu_1 = pL$, $\mu_2 = pL^2$, and $\Pi_C^{RB} = 0$ trivially, since $R_i - Lp_i = 0$ for

every group; substituting into $CS^{perceived}(\mu_1, \mu_2)$ and completing the square gives

$$CS_C^{RB,perceived} = \frac{Ap(1+h)}{2} - pL + \frac{pL^2}{2A(1+h)} = \frac{p[A(1+h) - L]^2}{2A(1+h)} = \frac{p\theta^2}{2A(1+h)} \quad (\text{A.84})$$

Under risk-based monopoly, $R_i = p_i\beta/2$ gives $\mu_1 = p\beta/2$, $\mu_2 = p\beta^2/4$ and substituting these gives

$$CS_M^{RB,perceived} = \frac{p\theta^2}{8A(1+h)} \quad (\text{A.85})$$

$$W_M^{RB,perceived} = \frac{3p\theta^2}{8A(1+h)} \quad (\text{A.86})$$

which is the standard monopoly-welfare fractional relationship under linear demand.

A.8.4 Welfare and Consumer Surplus Comparisons

Both comparisons below are Pool minus RB, so a positive sign means pooling raises the relevant measure and a negative sign means pooling lowers it.

Competition. Zero profit, $\Pi(R) = 0$, rearranges to $R^2q(\kappa) = R\beta - LpA(1+h)$; substituting this into $CS^{perceived}(R)$ eliminates its quadratic term at the zero-profit price R^{PC} , giving

$$CS_C^{Pool,perceived} = \frac{\theta}{2} \left[p - \frac{R^{PC}}{A(1+h)} \right] \quad (\text{A.87})$$

Since $\Pi_C^{Pool} = \Pi_C^{RB} = 0$, welfare equals consumer surplus in both cases, so subtracting $CS_C^{RB,perceived}$ above gives a closed form directly:

$$W_C^{Pool,perceived} - W_C^{RB,perceived} = CS_C^{Pool,perceived} - CS_C^{RB,perceived} = \frac{\theta(pL - R^{PC})}{2A(1+h)} \quad (\text{A.88})$$

As with the objective comparison in Section 5.2, this is determined by $pL - R^{PC}$, so total subjective welfare is lower under pooling than risk-based pricing whenever risk is heterogeneous.

Monopoly. Under pooled monopoly, $\mu_1 = \beta/[2q(\kappa)]$ and $\mu_2 = \beta^2/[4q(\kappa)]$. Substituting these into $CS^{perceived}(\mu_1, \mu_2)$ and $W^{perceived}(\mu_1, \mu_2)$, and subtracting the risk-based benchmarks above, the consumer surplus difference factors as

$$CS_M^{Pool,perceived} - CS_M^{RB,perceived} = \frac{(pq(\kappa) - 1)\beta[3A(1+h) - L]}{8A(1+h)q(\kappa)} \quad (\text{A.89})$$

Since $3A(1+h) - L > A(1+h) - L > 0$ always (given $A(1+h) > L$, assumed throughout), this is strictly positive for all $\kappa > 0$: subjective consumer surplus alone rises from pooling under monopoly unconditionally.

The welfare difference collects into a term free of $q(\kappa)$ and a term proportional to $1/q(\kappa)$, which factor identically, but with $3A(1+h) - L$ replaced by the smaller quantity $A(1+h) - 3L$:

$$W_M^{Pool,perceived} - W_M^{RB,perceived} = \frac{(pq(\kappa) - 1) \beta [A(1+h) - 3L]}{8A(1+h) q(\kappa)} \quad (\text{A.90})$$

Unlike the consumer surplus result, this is only positive when $A(1+h) > 3L$; when $A(1+h) < 3L$, pooling raises consumer surplus but can reduce profits by more than this when $A(1+h) < 3L$.

B Numerical Methods for the Distributional Robustness Check

Figure 9 shows monopoly premiums with both p_i and ν_j distributed as truncated Normal distributions instead of the paper's main specification, $\delta_i \sim U[-\kappa, \kappa]$ and $\nu_j \sim U[\ell, h]$. The distributions are truncated at $p_i > 0$ and $\nu_j > -1$ respectively, since neither a negative objective risk nor a subjective probability below zero is economically meaningful. Both truncated Normal distributions are constructed so that their own post-truncation mean and variance exactly match the corresponding uniform case's mean and variance, for the same κ and h .

Solving for the “parent” Normal’s parameters. For a target mean and variance (μ, σ^2) and a one-sided truncation point a , the parent Normal parameters (μ_p, σ_p) are found by solving the two-equation system given by the standard truncated Normal moment formulas

$$\mu = \mu_p + \sigma_p \lambda(\alpha), \quad \sigma^2 = \sigma_p^2 [1 - \lambda(\alpha)(\lambda(\alpha) - \alpha)], \quad \alpha \equiv \frac{a - \mu_p}{\sigma_p}, \quad \lambda(\alpha) \equiv \frac{\phi(\alpha)}{1 - \Phi(\alpha)} \quad (\text{B.91})$$

where ϕ and Φ are the standard Normal density and distribution function and $\lambda(\alpha)$ is the inverse Mills ratio. This system is solved numerically via Newton's method with a numerical Jacobian, starting from the target mean and standard deviation as an initial guess. For risk, the target is $(\mu, \sigma^2) = (1, \kappa^2/3)$ (matching $\delta_i \sim U[-\kappa, \kappa]$'s own moments, applied to p_i/p) with $a = 0$; for belief, the target is $\left(\frac{1+h}{2}, \left[\frac{1+h}{\sqrt{12}}\right]^2\right)$ (matching $\nu_j \sim U[-1, h]$'s own moments) with $a = -1$.

Risk-side quadrature. Quadrature nodes and weights for the truncated Normal risk distribution are obtained via an inverse-CDF method: 100-point Gauss-Legendre nodes on $[-1, 1]$ are mapped onto the probability scale $(0, 1)$, then passed through the truncated Normal's own quantile function, using the solved parent parameters.

Belief-side specification. Demand at a given premium R and risk realization p_i is the truncated Normal survival function of ν_j , evaluated at the implied belief cutoff $R/(Ap_i) - 1$: the untruncated

Normal survival function divided by the constant $1 - \Phi\left(\frac{-1-\mu_p}{\sigma_p}\right)$, the total probability mass retained above the truncation point, using belief's own solved parent parameters.

Profit maximization. For each (κ, h) , the monopolist's premium is found by univariate maximization of $\Pi(R) = R \sum_k w_k D_k(R) - L \sum_k w_k p_k D_k(R)$ over the risk quadrature nodes k , using a bounded Brent search method.

Grid. $\kappa \in [0, 0.8]$ (200 points) for each of $h \in \{1, 1.5, 2, \dots, 4.5\}$ (8 values). The upper bound of 0.8 is capped, as elsewhere in the paper, to keep every consumer's perceived loss probability below one; see the discussion in Section 4.