

## ADVERSE SELECTION IN AN INSURANCE MARKET WITH GOVERNMENT-GUARANTEED SUBSISTENCE LEVELS

Bum J. Kim  
Harris Schlesinger

### ABSTRACT

We consider a competitive insurance market with adverse selection. Unlike the standard models, we assume that individuals receive the benefit of some type of potential government assistance that guarantees them a minimum level of wealth. For example, this assistance might be some type of government-sponsored relief program, or it might simply be some type of limited liability afforded via bankruptcy laws. Government assistance is calculated *ex post* of any insurance benefits. This alters the individuals' demand for insurance coverage. In turn, this affects the equilibria in various insurance models of markets with adverse selection.

### INTRODUCTION

Governments often help in protecting their citizenry against risks. This may take many forms. For example, governments might offer public insurance; they might reinsure particularly problematic catastrophic risks; or they might provide for low-interest loans following a severe loss. Oftentimes, this relief takes the form of guaranteeing some minimum level of wealth. Poorer families suffering a loss might receive direct transfer payments to bring them up to some predetermined minimum wealth level. Or consider bankruptcy laws that allow for one to shield some prespecified level of wealth against creditors.

In this article, we consider only the case in which government assistance takes the form of guaranteeing some minimum wealth level, such as via direct government transfers following a loss. Our focus is not on the merits of having this type of assistance program in place; rather we examine its effects within an insurance market subject to adverse selection. In particular, we already know that adverse selection itself imposes some welfare costs, since any type of efficiency that is obtained must be "second-best" in nature. In this article, we examine how such second-best insurance contracting is

---

Bum J. Kim is an Assistant Professor at California State University at Bakersfield. Harris Schlesinger is the Samford Chain of Insurance at the University of Alabama and Adjunct Professor at the University of Konstanz. The author can be contacted by e-mail: hschlesi@cba.ua.edu. The authors are grateful to an anonymous referee for comments that helped to improve the article.

affected by the existence of the government assistance. We pay particular attention to how the welfare costs (often referred to as “signaling costs”) associated with the adverse selection are affected.

Except for noneconomic reasons, such as personal pride, government assistance can act as a substitute for market-based insurance: the possibility of government assistance might lower the demand for insurance. Or perhaps not. Consider the simple case of a two-state loss versus no-loss model. Since government subsidy programs are typically written to be “excess” of insurance coverage, government benefits are calculated only after all insurance indemnities have been paid out. For instance, purchasing a level of insurance that would leave one at a wealth level equal to the government-guaranteed minimum in the loss state would be totally redundant. One could receive the same wealth level in the loss state via government assistance, without the purchase of insurance. Indeed, one would be better off with government assistance, since there would not be any insurance premium to pay in the no-loss state of the world. The government essentially provides “free insurance.”

On the other hand, if insurance prices were actuarially fair, one would want to purchase full coverage insurance in the absence of any governmental assistance programs. Of course, this “fair premium” is based upon the full amount of the loss. By paying the premium for full coverage, an individual must give up the value of the government assistance. This creates a type of fixed cost for purchasing full insurance. Put differently, the individual would need to weigh the relative benefit of receiving the minimal governmental level of “insurance” for a zero premium versus paying a market-based insurance premium in return for a full-coverage insurance contract.

Now suppose that, for some reason, insurance coverage available via the marketplace were limited. In this case, the tradeoff between government assistance and market insurance would be different. For example, in the extreme, if the level of insurance available in the marketplace left the individual with no more wealth in the loss state than he or she would have via government assistance, there would be a zero demand for market insurance. More generally, since the marginal price of insurance is still fair, the individual either (1) will buy as much insurance as is available, or (2) will not purchase any insurance at all.

This point was made theoretically by Shavell (1986), in a liability context. The main focus of Shavell’s article is to consider the effects of limited liability on the optimal level of care taken by a potential injurer. A later article by Kaplow (1991), considers the effects of government relief on optimal care and optimal risk allocation. In this article, we assume that there are no moral hazard issues and that loss probabilities are fixed. Instead, we examine the effects of government relief or bankruptcy shields on equilibria in a market characterized by adverse selection.

If we consider a competitive insurance market in the presence of adverse selection, the value of such governmental assistance is likely to differ among the various risk types in the population. This is due to two reasons in particular. First, it may be the case that each risk type would self-select a different level of market insurance coverage in an equilibrium. Thus, each type must compare its own unique “market contract” with the option of forgoing insurance and accepting government assistance. Second, the different loss probabilities affect the relative valuation of the market insurance contracts offered *vis-à-vis* the alternative of government assistance.

In this article, we examine how government involvement in providing subsistence-level wealth affects equilibria in adverse-selection economies. We show the effects upon both the nature and the existence of equilibrium in three particular cases: the classic Nash equilibrium model of Rothschild and Stiglitz (1976), the potential pooling equilibrium of Wilson (1977), and the cross-subsidization equilibrium of Miyazaki (1977) and Spence (1978).

We start in the next section with a simple model of insurance demand in a world with competitive insurance markets and government assistance programs of the type mentioned above. We next examine the effects of such programs on the various models of adverse selection, paying particular attention to the welfare effects of the various participants.

### GOVERNMENT ASSISTANCE UNDER COMPLETE INFORMATION

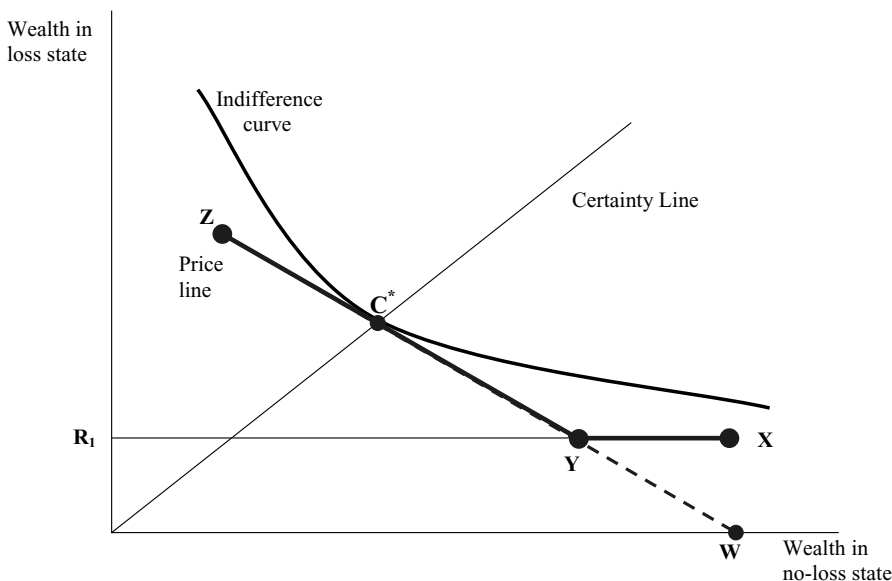
In this section we set up a simple model of government-guaranteed subsistence levels and their affect upon insurance demand. Our model is essentially an adaptation of Shavell (1986).

Consider a risk-averse expected-utility-maximizing individual with an initial wealth level  $W$ . The individual is exposed to a loss of a fixed size  $l \leq W$ , which occurs with probability  $p$ ,  $0 < p < 1$ . To simplify the exposition, we assume that the loss is total, if it occurs, i.e.,  $l = W$ . We assume that  $p$  is fixed, so that there are no moral hazard issues. An insurance contract is specified as an ordered pair  $(\alpha, \beta)$ , where  $\alpha$  denotes the premium to be paid in the no-loss state in return for a net benefit of  $\beta$  in the loss state. We assume that insurance markets are competitive with the premium for any insurance contract set at an actuarially fair price. If we let  $I$  denote the gross indemnity, then  $\beta = I - \alpha$ . In other words,  $\beta$  is simply the indemnity benefit net of the premium paid. Actuarially fair pricing implies that the premium equals the expected gross indemnity:  $\alpha = pI$ , or equivalently  $\alpha = p\beta/(1 - p)$ . If the individual is allowed to choose any level of benefit  $\beta$ , under the fair-pricing assumption, it is well known since Mossin (1968) that the individual will choose full coverage:  $\alpha = pW$ ,  $\beta = (1 - p)W$ .

Now suppose that the government establishes a minimum level of consumption, say  $R_1 > 0$ . If wealth falls below  $R_1$ , the government will provide transfer payments to the individual to make up the difference. This is illustrated via a standard state claims diagram in Figure 1. The individual's initial state claim is at  $(W, 0)$ . However, since the individual's state-2 wealth of zero falls below the guaranteed minimum of  $R_1$ , the government will provide the individual with a transfer payment of  $R_1$  to make up the difference. Thus, the individual's *de facto* initial state claim is  $(W, R_1)$ , which is depicted as point X in the diagram. For example, suppose the individual has an initial wealth of \$50,000 and might lose it all with a probability of  $p = 0.2$ . Say the government will ensure a wealth level of \$10,000. Thus, if a total loss occurs, the government provides a transfer payment of \$10,000.

Consider the purchase of insurance under these conditions. We assume that government assistance is only provided after all insurance indemnification has been paid out. Suppose insurance prices are actuarially fair. The individual can have essentially \$10,000 worth of "insurance coverage" with a zero premium, via the government.

**FIGURE 1**  
Government Relief and Effective Insurance Prices



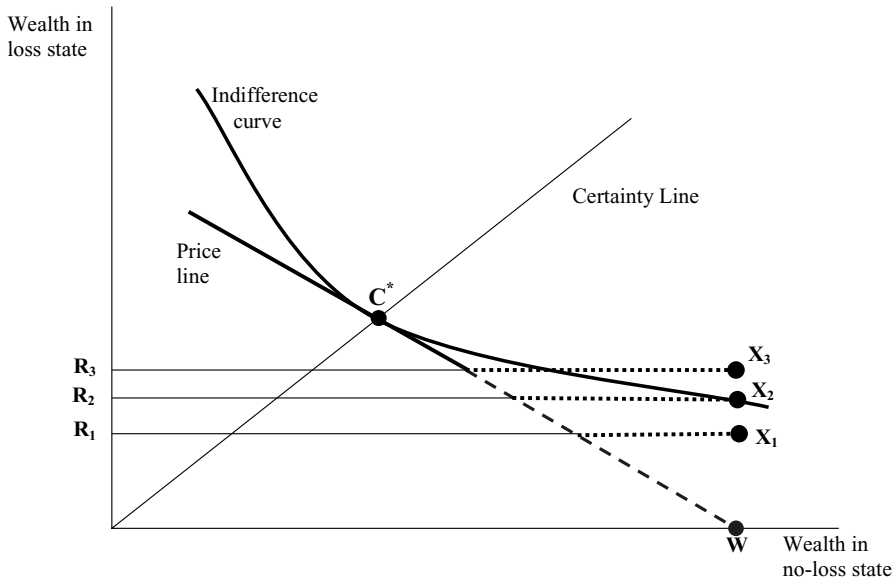
But if the individual desires a wealth level of, say, \$10,001 in the loss state, she cannot simply buy the extra \$1 in net coverage. She must purchase \$12,501.25 worth of insurance at the fair price, which amounts to an insurance premium of \$2500.25, for a net wealth of  $\beta = \$10,001$  if in the loss state. Thus, the extra net \$1 benefit is very expensive. Obviously, in this circumstance, the individual would rather have the free government assistance. In essence, we can view the opportunity cost of the government assistance as a type of “fixed cost” for obtaining marginal insurance coverage. But since marginal pricing of insurance is fair, if the individual decides to purchase insurance, she will opt for full coverage.<sup>1</sup>

In Figure 1, this can be seen by noting that the line **WYZ** denotes the fair-price line for insurance. However, purchases for insurance with  $\beta < R_1$  will be state dominated by the government assistance. Indeed, the true decision of the individual will be to choose the best state claim along the piecewise-linear, kinked line **XYZ**. Given risk aversion and fair pricing, this leads to a simple comparison of wealth with full coverage (at **C\***) or wealth with zero coverage (at **X**). As drawn in Figure 1, we see that full coverage will yield a higher level of expected utility than zero coverage.

Of course, the choice of full coverage versus zero coverage will depend upon the individual circumstances. For example, in Figure 2, we consider two higher levels of government-guaranteed wealth  $R_3 > R_2 > R_1$ . As shown in Figure 2, we see that at government-guaranteed wealth level,  $R_3$ , the individual prefers no insurance to full insurance (state claim **X<sub>3</sub>** is preferred to **C\***). Similarly, we see that at government-guaranteed wealth level,  $R_2$ , the individual is indifferent between full coverage and no

<sup>1</sup> This follows easily from Mossin (1968).

**FIGURE 2**  
Government Relief Versus Insurance



coverage. Indeed, it is obvious that we obtain a type of “bang-bang” solution whereby there exists a unique cut-off level of government-guaranteed wealth,  $R^*$ , such that zero insurance is optimal if and only if the guaranteed wealth  $R$  is more than  $R^*$ .<sup>2</sup> Of course,  $R^*$  is unique to the individual and need not be the same for everyone. For example, it is easy to show that, *ceteris paribus*, a higher degree of absolute risk aversion leads to a higher cut-off level of  $R^*$ .

### ADVERSE SELECTION AND SEPARATING EQUILIBRIA

Here, we consider the adverse-selection model of Rothschild and Stiglitz (1976). There are two types of individuals within one risk classification, who differ only in their probabilities of a loss. The “good risks” have a probability of loss  $p_G$  that is lower than the corresponding loss probability of the “bad risks,”  $p_B$ ,  $0 < p_G < p_B < 1$ . Individuals have private information about their own type while insurers only know the distribution of the two types within the population, i.e., they know  $\lambda$ , the fraction of bad risks in the population.

Rothschild and Stiglitz consider a competitive insurance market in which consumers are offered a menu of contracts. Insurers are risk neutral and seek to maximize expected profit. Each insurance contract consists of an  $(\alpha, \beta)$  pair, denoting the premium and the net indemnity. Individuals are free to choose from among all of the contracts made available. An equilibrium set of contracts is characterized as follows:

<sup>2</sup> This was essentially shown by Shavell (1986), who examines liability for damages that might exceed one’s wealth. In Shavell’s model, initial wealth varies, whereas here, the level of government-guaranteed wealth (which could be a bankruptcy shield) varies while initial wealth stays fixed.

- E1 Every contract type in the equilibrium set expects to earn a zero profit.
- E2 There does not exist a contract that is not in the equilibrium set that would be able to earn a positive profit, on average, if added to the set of equilibrium contracts.

In order to rule out contracts with zero demand, we also assume that each contract type in the equilibrium set is purchased by some consumers. Condition E1 is an assumption that firms in competitive markets earn zero profit in the long run. Condition E2 is that the resulting equilibrium is a Nash noncooperative equilibrium.<sup>3</sup>

If an equilibrium exists, Rothschild and Stiglitz show that it must be a separating equilibrium designed as follows. The bad risks receive full insurance at a fair price for the bad risks:  $\alpha = p_B W$ ,  $\beta = (1 - p_B)W$ . The good risks receive a level of coverage at the good-risk fair price, but it is restricted by an incentive-compatibility constraint, whereby the bad risks must not prefer the good-risk contract to their own. In particular, we allow only a restricted amount of coverage,  $\alpha_G$ , to be purchased at the lower good-risk price. The level  $\alpha_G$  is chosen such that the bad risks are indifferent between the contingent wealth pair with full insurance at the bad-risk price,  $[(1 - p_B)W; (1 - p_B)W]$ , and the contingent wealth pair under restricted level of coverage at the good-risk price  $[W - \alpha_G, (1 - p_G)\alpha_G/p_G]$ . These two contingent wealth levels are depicted as **B** and **G** in Figure 3.<sup>4</sup>

The separating equilibrium consists of two contracts, yielding contingent allocations **G** and **B**. The good risks self-select **G** while the bad risks prefer **B**. It is interesting to note the welfare implications of the adverse selection here. In a market with complete information, the bad risks would have full insurance at allocation **B**, same as they do with the adverse selection. Thus, there is no loss or gain in welfare for the bad risks. The good risks, however, would be able to buy full insurance at a good-risk fair price under complete information and achieve allocation **H**. The good risks thus must bear a welfare loss by restricting their coverage to  $\alpha_G$ . This is often considered a "signaling cost" that must be borne in order to obtain the favorable good-risk price.

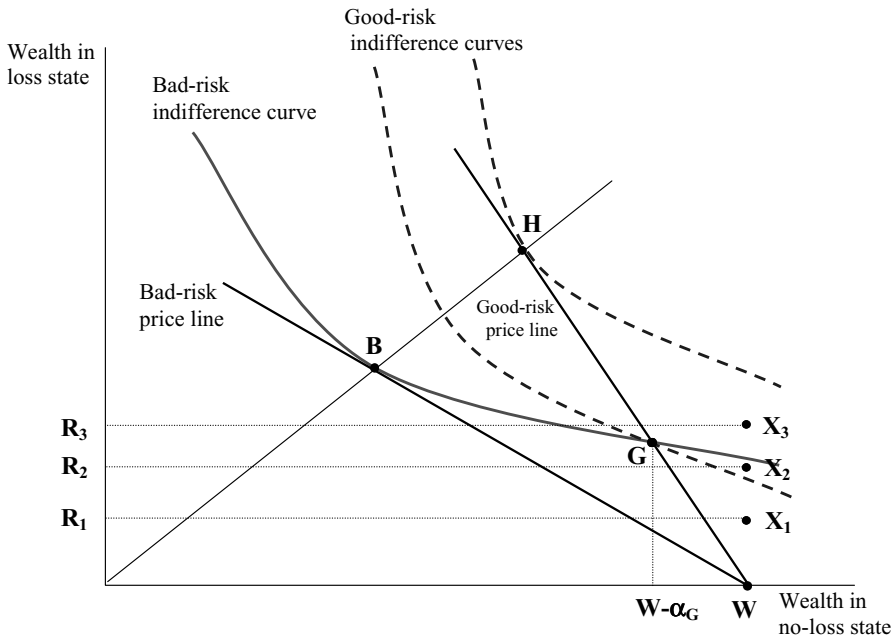
Let us consider now, how the existence of government-guaranteed minimum wealth affects this equilibrium setting. At very low levels of government guarantees, the equilibrium will not be affected at all. This is clearly the case with the government guarantee set at  $R_1$  in Figure 3.

At the other extreme, if the government-guaranteed minimum wealth is high enough, such as at level  $R_3$  in the figure, both types of individuals will prefer zero coverage to market insurance. In this setting, all individuals will have contingent wealth  $X_3$ . Note that, as drawn, the good risks still have a welfare loss due to signaling costs, since in

<sup>3</sup> Of course, insurers might have costs to be met in the real world, but they also have investment opportunities from premium funds, which do not manifest themselves in our static model. Thus, the zero-profit assumption is not too unrealistic. The equilibrium is Nash if we view contract offers as the strategy space for insurers. E2 then requires that strategies be mutual best responses.

<sup>4</sup> Although the bad risks are indifferent to these two contracts, we assume they opt for full coverage. In a more realistic discrete setting, we could assume that  $\beta_G$  is reduced by 1 cent, so that there is a strict preference by the bad risks.

**FIGURE 3**  
Rothschild–Stiglitz Separating Equilibrium



the absence of adverse selection they would achieve contingent wealth **H**, which is strictly preferred to  $X_3$ .

Of course, we cannot fully compare welfare levels with and without government assistance. This would require taking the source of government financing into account. To say that the signaling cost is lower under government subsistence level,  $R_3$ , is true; however, the total welfare effect on the good risks would need to consider any taxes paid in order to finance government relief. However, we can assume a world with government relief in place and then consider the welfare effects (signaling costs) of the adverse selection, compared to a world with no government assistance.<sup>5</sup> Under full information, each type chooses either full coverage at its own fair price, or it chooses zero coverage and accepts government assistance. Thus, under government relief level,  $R_3$ , adverse selection has a lower welfare loss (since the bad risks choose  $X_3$  with or without information asymmetry and the good-risk welfare loss at  $X_3$  is less than at **G**). At the lower level of assistance,  $R_1$ , the welfare loss from information asymmetry is the same as in the world without government relief.

Finally, since the good-risk indifference curve is always steeper than the bad-risk indifference curve through the same contingent wealth, there will be some levels

<sup>5</sup> This comparison assumes that any income effects, due to higher taxes in the world with government relief, are negligible. If we tell a story where government guarantees are via bankruptcy shields, established by law, then the minimum subsistence-level costs would be financed by unpaid claims due to injured third parties, as is the case in Shavell (1986).

of government guarantees for which the good risks prefer no insurance, while the bad risks still prefer full coverage. This is true at government guarantee of income level,  $R_2$ , in Figure 3, for example. Under  $R_2$ , a separating equilibrium offers only one insurance contract, full insurance at the bad-risk price. Only the bad risks will purchase this contract, whereas the good risks will purchase no coverage and opt for government relief.

If we consider the world with government relief level,  $R_2$ , in place, the bad risks receive **B** with or without the adverse selection, whereas the good risks obtain wealth claim  $X_2$  rather than **G**, representing the welfare cost of adverse selection, which is lower than the loss without the government relief.

### POOLING CONTRACTS AND EQUILIBRIUM

Rothschild and Stiglitz (1976) show that there cannot exist a pooling Nash equilibrium. However, Wilson (1977) extends the equilibrium class to one for which a pooling equilibrium is possible. In Wilson's "anticipatory" equilibrium, we replace equilibrium E2 with the following:

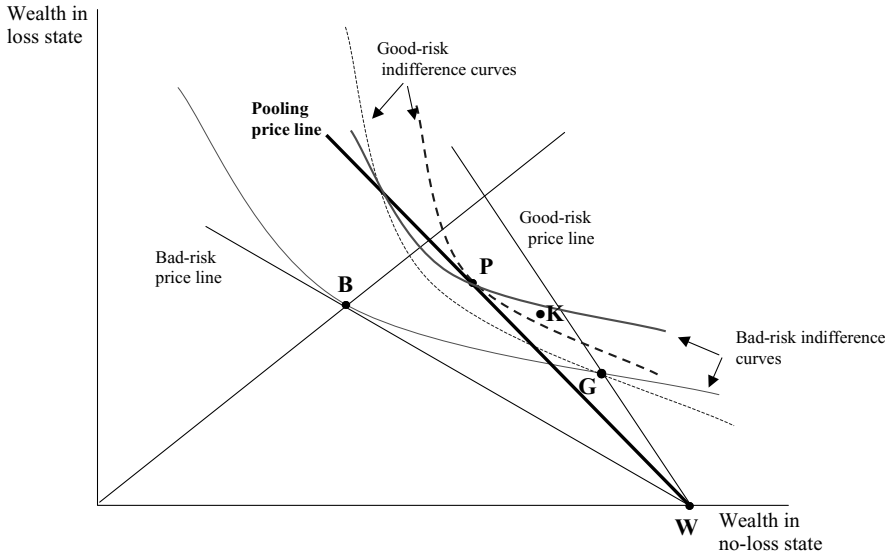
- E2' There does not exist a contract that, if added to the set of equilibrium contracts, would be able to earn a positive expected profit, even after nonprofitable contracts are withdrawn from the market.

To see the impact of Wilson's definition, let us first consider the possibility of a pooling contract in the Rothschild–Stiglitz model. Such a contract offers an  $(\alpha, \beta)$  contract pair with prices set such that  $\alpha = p_\lambda \beta / (1 - p_\lambda)$ , where  $p_\lambda = \lambda p_B + (1 - \lambda)p_G$  is the "fair pooling price" per dollar of indemnity. If everyone purchases the pooling contract, then any contract satisfying this price relationship will break even, on average. The contract earns a profit on the good risks, and generates a loss on the bad risks, but breaks even on average.

The pooling price line is drawn in Figure 4. It will lie somewhere between the good-risk and the bad-risk price lines, depending on the relative number of bad risks,  $\lambda$ . If  $\lambda$  is high enough, so that the pooling price line lies close to the bad-risk price line, then the good risks will all prefer their separating contract  $(\alpha_G, \beta_G)$ , at wealth level **G**, to all possible pooling contracts. However, if  $\lambda$  is low enough, such as drawn in Figure 4, there will be some pooling contracts that are preferred by the good risks to  $(\alpha_G, \beta_G)$ . For example, consider the insurance contract  $(\alpha_P, \beta_P)$  that leaves everyone with contingent wealth **P**. This contract would earn zero profit on average if both types purchase it. Contingent wealth **P** is preferred by the good risks to **G** and by the bad risks to **B**. Hence, everyone would indeed purchase the pooling contract  $(\alpha_P, \beta_P)$  if it were added to the set of separating contracts. Indeed, if we reduce  $\beta$  by some small amount  $\varepsilon$ , then the contract  $(\alpha_P, \beta_P - \varepsilon)$  would also be preferred by everyone, and this contract would earn a positive profit on average, thus negating the equilibrium properties of the separating contracts.

Of course, a positive profit could not persist in equilibrium, but why not an equilibrium pooling contract at  $(\alpha_P, \beta_P)$ ? Since the good-risk indifference curve through **P** must be

**FIGURE 4**  
Pooling Contracts



steeper than the bad-risk indifference curve, we can always find some new contract, such as  $(\alpha_K, \beta_K)$  yielding wealth level **K** in Figure 4, satisfying the following properties:

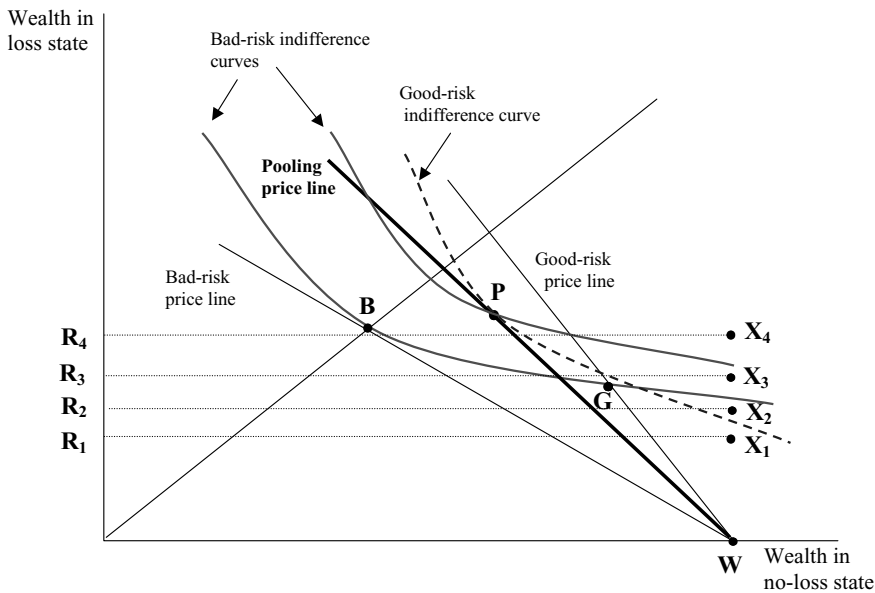
- A The good risks prefer contract  $(\alpha_K, \beta_K)$ .
- B The bad risks prefer contract  $(\alpha_P, \beta_P)$ .
- C Contract  $(\alpha_K, \beta_K)$  earns a profit, on average, if only the good risks buy it.

It is clear that contract  $(\alpha_P, \beta_P)$  cannot satisfy equilibrium condition E2, hence is not a Rothschild–Stiglitz equilibrium. Indeed, the same arguments were used by Rothschild and Stiglitz to show that no pooling contract will be a Nash equilibrium.

Using Wilson’s property E2’ in our definition of equilibrium, changes matters, however. Under Wilson’s equilibrium, if  $\lambda$  is high enough, the pair of Rothschild–Stiglitz separating contracts will remain the equilibrium. However, in a situation where  $\lambda$  is low enough, such as depicted in Figure 4, there will exist a pooling equilibrium. Hence, there always exists an equilibrium in Wilson’s model. In particular, the zero-profit pooling contract that is the most preferred by the good risks will be the equilibrium pooling contract. This is precisely the contract  $(\alpha_P, \beta_P)$  in Figure 4, yielding contingent wealth **P**. The reason for this is clear. Any profits will be driven out by long-run competition. And, any other zero-profit pooling contract, if offered concurrent with  $(\alpha_P, \beta_P)$  would attract no buyers among the good risks. Thus, such a contract would lose money, on average, if only the bad risks purchase it.

The reason that  $(\alpha_P, \beta_P)$  can support an equilibrium under Wilson’s definition is that a contract such as  $(\alpha_K, \beta_K)$  will not be offered. This is because the addition of contract

**FIGURE 5**  
Pooling Equilibrium and Government Relief



$(\alpha_K, \beta_K)$  initially would attract only the good risks, leaving the bad risks as the only ones purchasing  $(\alpha_P, \beta_P)$ . In this event, the contract  $(\alpha_P, \beta_P)$  would lose money and hence be withdrawn from the market. However, if this happens, then  $(\alpha_K, \beta_K)$  would attract both good risks and bad risks, and would thus also lose money. As a result, the contract  $(\alpha_P, \beta_P)$  is a pooling equilibrium.

So how does the existence of a government-guaranteed minimum wealth affect potential pooling-equilibrium contracts? Obviously, if the level of government relief is very low, such as at level  $R_1$  in Figure 5, then it will have no effect upon a pooling equilibrium. Contingent claim  $P$  dominates  $X_1$  for both risk types in Figure 5. Similarly, if the government-guaranteed wealth level is sufficiently high, such as at level  $R_4$  in Figure 5, both types of individuals will prefer to accept the government relief and purchase no insurance coverage. In Figure 5, contingent claim  $P$  is dominated by  $X_4$  for both risk types.

The interesting cases occur when the level of government relief is high enough to attract just one of the two types of risks. Since good-risk indifference curves are everywhere steeper, this can only occur by attracting the good risks. Such is the case at the government-guaranteed wealth level  $R_3$  in Figure 5. This would yield a contingent-wealth claim of  $X_3$  with no insurance versus  $P$  with the pooling contract  $(\alpha_P, \beta_P)$ . In this setting, the bad risks would prefer  $(\alpha_P, \beta_P)$ , whereas the good risks are better off with no insurance. In a Rothschild–Stiglitz-type setting, this would preclude  $(\alpha_P, \beta_P)$  from being an equilibrium, since only the bad risks purchase it and, hence, it loses money. The same is true under Wilson’s definition of equilibrium. In both settings, equilibrium would entail both types of individuals purchasing zero coverage, since  $X_3$  is preferred to full coverage by the bad risks.

Another interesting possibility can occur in some circumstances, as is illustrated in Figure 5 at government-guaranteed wealth level,  $R_2$ .<sup>6</sup> First of all, let us recall that in the absence of any government relief, there would be no Rothschild–Stiglitz equilibrium in the setting for Figure 5. If we use Wilson’s concept of equilibrium, then we would obtain the pooling equilibrium contract  $(\alpha_P, \beta_P)$ . Now, at the level of minimum wealth,  $R_2$ , the good risks prefer no insurance at contingent wealth  $X_2$ , while the bad risks prefer the pooling contract  $(\alpha_P, \beta_P)$  at wealth  $P$ . As a result, if the pooling contract is offered, it will lose money, since only bad risks will purchase it. Thus  $(\alpha_P, \beta_P)$  is not viable. However, note that the bad risks prefer full coverage at the bad-risk fair price to zero coverage. As a result, we will obtain a separating equilibrium, in which the bad risks purchase full coverage at contingent wealth  $B$  and the good risks purchase no coverage at contingent wealth level,  $X_2$ . This is true for either the Wilson or the Rothschild–Stiglitz definition of equilibrium. In other words, we go from either no equilibrium (in the Rothschild–Stiglitz sense) or a pooling equilibrium (in Wilson’s sense) to a separating equilibrium when the government guarantee is set at  $R_2$ .

To consider the welfare effects due to the informational asymmetry, at relief level  $R_1$ , we obtain a Wilson pooling equilibrium and the welfare effect is the same as without government intervention. At  $R_4$  the welfare effect depends on whether or not the good risks prefer  $X_4$  to full insurance at the good-risk price. If they do, there is no welfare loss. Otherwise, the loss is simply the difference in good-risk utility between  $X_4$  and full coverage. At  $R_3$ , the bad risks have  $X_3$  with or without the informational asymmetry, but the good risks also receive contingent wealth  $X_3$ , which might or might not be preferred to full coverage. At  $R_2$ , the good risks do not purchase insurance, which again might or might not be preferred to the full insurance available in the full-information case. The bad risks will have full coverage at  $B$ , same as in the full-information case.

If we compare these welfare losses to those in the no-government-relief world, obviously  $R_4$  leads to lower welfare loss, since everyone is better off at  $X_4$  than at  $P$ . But at  $R_2$  and  $R_3$  we have a tradeoff. In both cases the good risks are better off than at  $P$ , whereas the bad risks prefer  $P$  to either  $X_3$  (in the case where the government-guaranteed minimum wealth is  $R_3$ ) or to  $B$  (in the case where the government-guaranteed minimum wealth is  $R_2$ ). Thus, any general claim about higher or lower welfare costs due to the informational asymmetry depends upon the welfare criterion used.

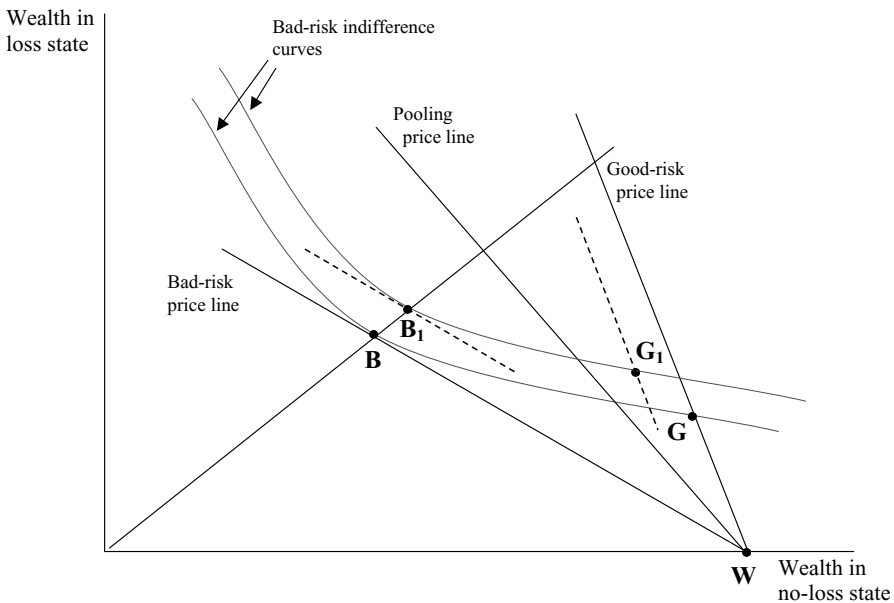
### SUBSIDIZING CONTRACTS AND EQUILIBRIUM

The Rothschild–Stiglitz definition of equilibrium included the long-run-competition condition E1, that every type of contract offered earns an expected profit of zero. However, an alternative assumption is the following:

E1’: Every insurer’s set of equilibrium contracts earns a zero profit, on average.

<sup>6</sup> Situations such as  $R_1$ ,  $R_3$ , and  $R_4$  always exist. A situation as depicted at  $R_2$  might or might not exist depending upon where the good-risk indifference curve through  $P$  and the bad-risk indifference curve through  $B$  intersect. If this occurs at a higher no-loss wealth level than  $W$ , then a situation such as that at  $X_2$  will not be possible.

**FIGURE 6**  
Subsidizing Contracts



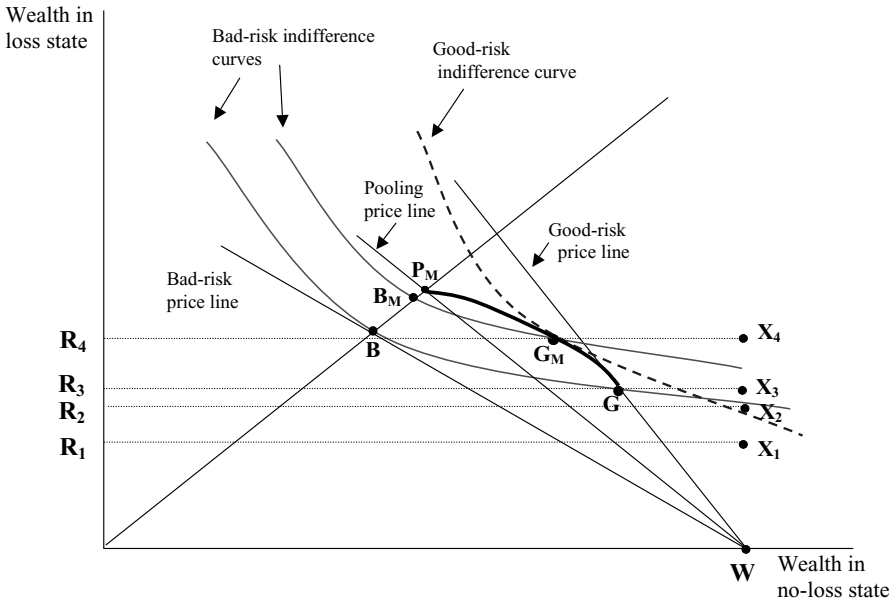
Condition E1' allows for some types of contracts to earn a positive profit, so long as other contracts offered expect to generate a loss. The key competitive assumption is that the menu of all contracts must break even, as opposed to assuming that every contract breaks even.

Miyazaki (1977) and Spence (1978) use this assumption, along with Wilson's anticipatory-equilibrium condition E2' to examine a class of subsidizing equilibria. To see how this works, consider the Rothschild–Stiglitz set of separating contracts: full insurance ( $\alpha_B, \beta_B$ ) for the bad risks with partial insurance ( $\alpha_G, \beta_G$ ) for the good risks. This provides the bad risks and the good risks with contingent wealth levels **B** and **G**, respectively, as illustrated in Figure 3 and again in Figure 6.

Now suppose that we allow the bad risks to receive full insurance at a premium priced with a subsidy of  $\sigma > 0$ . In other words, the bad risks receive the full-coverage contract ( $\alpha_B - \sigma, \beta_B + \sigma$ ). This is illustrated as yielding the contingent wealth **B<sub>1</sub>** in Figure 6. In order to finance this subsidy, the good risks are offered insurance with a fair marginal price, but with a "tax" of  $\tau$  per policy. The break-even condition E1' implies that we must have  $\tau = \lambda\sigma / (1 - \lambda)$ . The good-risk contract thus must include this tax by adding  $\tau$  to the premium in the no-loss state of nature and subtracting  $\tau$  from the net indemnity in the loss state.

The level of insurance coverage offered to the good risks is once again limited by the incentive-compatibility constraint of the bad risks. The good risks are offered as much coverage as possible, without providing an incentive for the bad risks to self-select the good-risk contract. This is illustrated as the contract leading to contingent wealth **G<sub>1</sub>** in Figure 6. Thus the pair of insurance contracts leading to the contingent wealth levels

**FIGURE 7**  
Miyazaki–Spence Equilibrium



$B_1$  and  $G_1$  in Figure 6 satisfies the zero-profit condition  $E1'$  and it is possibly preferred by both the good risks and the bad risks to the Rothschild–Stiglitz separating pair of contracts; the bad risks definitely prefer  $B_1$  to  $B$  and the good risks might prefer  $G_1$  to  $G$ . Might this subsidizing set of contracts be an equilibrium in the sense of Wilson?

Of course, we can always adjust the subsidy and tax to find more such subsidizing pairs. The set of all good-risk wealth allocations derived from these subsidizing contracts is illustrated as the curve  $GG_MP_M$  in Figure 7. For each contingent-wealth allocation for the good risks along this locus, the corresponding bad-risk wealth allocation is determined by the bad-risk indifference curve through this allocation: the bad risks receive full insurance yielding this level of expected utility. For example, at allocation  $G_M$  for the good risks, the bad risks would receive full insurance at contingent-wealth allocation  $B_M$ . As another example, the wealth allocation at  $P_M$  is for full insurance at a pooling price. Such an insurance contract, when purchased by both types, is also a zero-profit-subsidizing “pair” of contracts.

As shown in Figure 7,  $G_M$  yields the most preferred level of contingent wealth available to the good risks from among all subsidizing contracts. Let  $(\alpha_G^M, \beta_G^M)$  denote the contract pair associated with wealth  $G_M$ . Similarly, let  $(\alpha_B^M, \beta_B^M)$  denote the corresponding full-insurance contract associated with contingent wealth  $B_M$ .

If another subsidizing pair of contracts  $[(\alpha_G^s, \beta_G^s); (\alpha_B^s, \beta_B^s)]$  was offered it could not be the equilibrium. The reasoning is as follows. Since  $(\alpha_B^M, \beta_B^M)$  is the most preferred contract for the good risks, any other contingent wealth along the curve  $GG_MP_M$  in Figure 7 yields less utility for the good risks. Thus, a new contract could be offered with net benefit  $\beta_G^M$  for the good risks, but with a premium slightly higher than  $\alpha_G^M$ . This

new contract,  $(\alpha_G^M + \varepsilon, \beta_G^M)$ , would be preferred to  $(\alpha_G^s, \beta_G^s)$  by the good risks. This contract could be offered simultaneously with contract  $(\alpha_B^M, \beta_B^M)$  for the bad risks. Thus, the contract pair  $[(\alpha_G^M + \varepsilon, \beta_G^M); (\alpha_B^M, \beta_B^M)]$  would earn a slight profit if both types of contracts were purchased. Whether the bad risks are better off with  $(\alpha_B^s, \beta_B^s)$  or with  $(\alpha_B^M, \beta_B^M)$  is irrelevant here, since the contract pair  $[(\alpha_G^s, \beta_G^s); (\alpha_B^s, \beta_B^s)]$  would lose money for an insurer if no good risks purchase  $(\alpha_G^s, \beta_G^s)$ . Hence, under condition E2' for a Wilson-type of equilibrium, the contract  $(\alpha_B^s, \beta_B^s)$  would be withdrawn from the market. Consequently, the contract pair  $[(\alpha_G^M, \beta_G^M); (\alpha_B^M, \beta_B^M)]$  is a Miyazaki–Spence equilibrium.<sup>7</sup>

The addition of government relief to the model has the same consequences as in the case of a Wilson pooling equilibrium. If the level of government-guaranteed minimum wealth is very low, such as  $R_1$  in Figure 7, then it has no effect on the private market equilibrium. If the level of government-guaranteed minimum wealth is very high, such as  $R_4$ , then both types will not purchase any insurance.

Now suppose that the level of government-guaranteed minimum wealth is at a level that makes the no-insurance option preferred by the good-risk type only. Since under a Wilson equilibrium the bad-risk contract  $(\alpha_B^M, \beta_B^M)$  would be withdrawn from the market, the bad risks will have to receive some other alternative. In this case we have two possible outcomes. If the guarantee level is low enough, such as at  $R_2$ , the bad risks will be offered full insurance at the bad-risk fair price. In this case the Miyazaki–Spence equilibrium will yield contingent wealth  $X_2$  for the good risks and  $B$  for the bad risks. However, if the guarantee level is a bit higher, say at  $R_3$ , the bad risks will prefer no insurance to full coverage. This implies a Miyazaki–Spence equilibrium in which no one buys market insurance and everyone has the same contingent-wealth claim  $X_3$ .

The welfare effects in the case of a Miyazaki–Spence are very similar to those in the case of a Wilson pooling equilibrium and not discussed in detail here. Once again, at  $R_2$ ,  $R_3$ , and  $R_4$ , we do not know without more information whether or not the good risks are better off with no insurance or with full insurance. Hence, the good-risk welfare might or might not be lower under adverse selection than under complete information. And once again the bad risks are worse off at either  $X_3$  or  $B$  than they are in the equilibrium, which here is at contingent wealth  $B_M$ .

## CONCLUDING REMARKS

We examined several models of insurance markets under adverse selection. Our point of departure is that we allow government subsidies to be considered as an alternative to market insurance. In all the models considered, a very low level of relief did not affect insurance markets at all, whereas a very high level of relief made all individuals

<sup>7</sup> Actually, there is no guarantee that the curve  $CG_MP_M$  will be concave. If it is not, we might find that the highest level of expected utility for the good risks is achievable via more than one good-risk contract. In such a case, the Miyazaki–Spence equilibrium contract is the one with the highest level of coverage, and hence highest tax, for the good risks. If this were not the case, some good risks might buy one of contracts with a lower subsidy. The bad risks, on the other hand, all will be best off buying the full-coverage contract with the highest subsidy. But this contract will lose money if a fraction less than  $1 - \lambda$  of the good risks pay the highest tax.

prefer government assistance to market insurance. In a Rothschild and Stiglitz (1976) setting, government relief might alter the set of separating-equilibrium contracts. In particular, we might find that the bad risks purchase full coverage while the good risks decide not to buy any insurance. This latter type of separation also might lead to an equilibrium in circumstances for which a Rothschild–Stiglitz type of equilibrium (i.e., Nash equilibrium) does not exist in the absence of government intervention.

In the case of Wilson's (1977) anticipatory equilibrium, we may find the same results as above. On the other hand, we might have a market with a Wilson pooling equilibrium, in which the equilibrium changes from a pooling to a separating type of equilibrium after government assistance is offered in the economy. Similarly, government assistance can alter the type of separation that would occur in a model similar to that of Miyazaki (1977) and Spence (1978).

In all of our analyses, we did not consider the source of financing for government assistance. In a large economy, with many public-goods projects, we can approximate this effect on the insureds as a deadweight loss on each individual. Of course, unless we use some severely strong criterion, such as the Pareto criterion, we need to balance off gains and losses in welfare to see the total value of such government assistance to society. However, any reduction in signaling costs should be a part of any such evaluation. Similarly, these signaling costs need to be included if we wish to compare alternative government assistance programs.

We also assumed that the level of wealth to be guaranteed as the government-guaranteed subsistence level was public knowledge. More realistically, individuals will only have some estimate of the level of government assistance that might be available or even whether such relief is forthcoming at all. In our model, this type of uncertainty would render the level of the guaranteed minimum wealth,  $\bar{R}$ , to be random, which in turn would be less valuable to risk-averse individuals. The effect of this type of randomness obviously will have an effect on the purchase of insurance, but this complicates the simplicity of our two-state model and is left to future research.

## REFERENCES

- Kaplow, L., 1991, Incentives and Government Relief for Risk, *Journal of Risk and Uncertainty*, 4: 167-175.
- Miyazaki, H., 1977, The Rat Race and Internal Labor Markets, *Bell Journal of Economics*, 8: 394-418.
- Mossin, J., 1968, Aspects of Rational Insurance Purchasing, *Journal of Political Economy*, 79: 553-568.
- Rothschild, M., and J. Stiglitz, 1976, Equilibrium in Competitive Insurance Markets: An Essay on the Economics of Imperfect Information, *Quarterly Journal of Economics*, 90: 629-650.
- Shavell, S., 1986, The Judgment Proof Problem, *International Review of Law and Economics*, 6: 45-58.
- Spence, M., 1978, Product Differentiation and Performance in Insurance Markets, *Journal of Public Economics*, 10: 427-447.
- Wilson, C., 1977, A Model of Insurance Markets With Incomplete Information, *Journal of Economic Theory*, 12: 167-207.

Copyright of Journal of Risk & Insurance is the property of Blackwell Publishing Limited and its content may not be copied or emailed to multiple sites or posted to a listserv without the copyright holder's express written permission. However, users may print, download, or email articles for individual use.