



CIRRELT

Centre interuniversitaire de recherche
sur les réseaux d'entreprise, la logistique et le transport

Interuniversity Research Centre
on Enterprise Networks, Logistics and Transportation

Economic Effects of Risk Classification Bans

Georges Dionne
Casey G. Rothschild

June 2014

CIRRELT-2014-30

Bureaux de Montréal :
Université de Montréal
Pavillon André-Aisenstadt
C.P. 6128, succursale Centre-ville
Montréal (Québec)
Canada H3C 3J7
Téléphone : 514 343-7575
Télécopie : 514 343-7121

Bureaux de Québec :
Université Laval
Pavillon Palasis-Prince
2325, de la Terrasse, bureau 2642
Québec (Québec)
Canada G1V 0A6
Téléphone : 418 656-2073
Télécopie : 418 656-2624

www.cirrelt.ca

Economic Effects of Risk Classification Bans

Georges Dionne^{1,*}, Casey G. Rothschild²

¹ Interuniversity Research Centre on Enterprise Networks, Logistics and Transportation (CIRRELT) and Department of Finance, HEC Montréal, 3000, Côte-Sainte-Catherine, Montréal, Canada H3T 2A7

² Wellesley College, Pendleton Hall East, Room 414, 106 College Street, Wellesley, MA 02481, USA

Abstract. Risk classification refers to the use of observable characteristics by insurers to group individuals with similar expected claims, to compute the corresponding premiums, and thereby to reduce asymmetric information. Permitting risk classification may reduce informational asymmetry-induced adverse selection and improve insurance market efficiency. It may also have undesirable equity consequences and undermine the implicit insurance against reclassification risk which legislated restrictions on risk classification could provide. We use a canonical insurance market screening model to survey and to extend the risk classification literature. We provide a unified framework for analyzing the economic consequences of legalized vs. banned risk classification, both in static-information environments and in environments in which additional information can be learned, by either side of the market, through potentially costly tests.

Keywords. Adverse selection, classification risk, classification bans, equity, efficiency.

Results and views expressed in this publication are the sole responsibility of the authors and do not necessarily reflect those of CIRRELT.

Les résultats et opinions contenus dans cette publication ne reflètent pas nécessairement la position du CIRRELT et n'engagent pas sa responsabilité.

* Corresponding author: Georges.Dionne@cirrelt.ca

Dépôt légal – Bibliothèque et Archives nationales du Québec
Bibliothèque et Archives Canada, 2014

© Dionne, Rothschild and CIRRELT, 2014

1. Introduction

Risk classification refers to the use of observable characteristics such as gender, race, behavior, or the outcome of genetic tests to price or structure insurance policies. Risk classification helps insurers to group individual risks with similar expected costs, to compute the corresponding insurance premiums, and to reduce adverse selection (and, potentially, moral hazard). Only risk characteristics correlated with expected claim costs are directly useful for underwriting.¹ Information on individual risk is seldom used to determine individual participation in employer- or government-sponsored insurance plans, but it is often observed in voluntary plans, where it serves to define accessibility, to classify policyholders into homogenous risk classes, and to set the premiums of each risk class. In health insurance, for example, premiums are commonly determined by age, sex, smoking behavior, parent's health history information, current medical conditions (high cholesterol, diabetes, etc.), and medical histories, particularly of older clients; information about lifestyle, diet, and exercise, though less commonplace, is also used for premium setting.

Risk classification may also take into account advances in diagnostics and treatments. These tests provide potentially useful and treatment-relevant information; they also raise two closely related concerns, particularly insofar as they reveal medical conditions that are exogenous to the individuals. First, there are the basic equity concerns that arise when individuals, through no fault of their own, find their access to affordable insurance curtailed by the outcome of such a test. Second, insofar as there is no well-developed market for insurance *against the outcome of the test*, these tests cause inefficiencies associated with the introduction of an additional *classification risk* (the risk of being re-classified and paying a higher premium or facing reduced access to insurance). These concerns have played a major role in the regulatory imposition of constraints on the use of certain types of information for insurance classification and pricing.² These concerns have been a particularly important driver of regulatory

¹ In non-competitive environments, characteristics not directly related to expected claims loss may also be employed, e.g., to identify individuals with higher willingness to pay.

² The Patient Protection and Affordable Care Act (ACA), passed in the U.S. in 2010, for example, entirely prohibits the use of individual health status in pricing insurance. Indeed, the Act expressly forbids pricing on any characteristics other than age, family status, geographic or rating area, and tobacco use and even restricts the use of these characteristics (Baker, 2011; Harrington, 2010a, 2010b). See also Morrissey (2013) for a general overview of the health insurance system in the United States.

restrictions on the use of genetic tests in insurance pricing—and, indeed, restrictions on the development and use of genetic tests more generally.³

In fact, restrictions on the use of risk classification are pervasive—and are motivated, at least in part on similar concerns for social equity and “classification risk” minimization. For example, gender-based risk classification has been prohibited in the European Union since late 2012.

Whether based on new diagnostic tests or pre-existing information such as gender or race, concerns about social equity and classification risk minimization may provide a strong *a priori* case for limiting the use of risk classification in insurance markets. As we discuss extensively in this survey, however, in settings with market based insurance provision, risk classification is also generally associated with increased efficiency.

To illustrate, observe that risk-pooling arising from legal restrictions on risk classification may lead to a situation in which lower-risk individuals are charged higher than actuarially fair premiums and higher-risk individuals are charged lower than actuarially fair premiums. While these financial inequities (may) reduce classification risk (and/or improve social equity), the higher-than-fair premiums for lower-risk individuals may cause them to forgo insurance entirely, particularly when the proportion of high-risk individuals is large. This reduced pool of insured individuals reflects a decrease in the efficiency of the insurance market. These negative efficiency consequences of limits on risk classification can, in principle, be quite severe: exit of low risks can lead to a “death spiral” of rising premiums and lower-risk exit that ends up unraveling and destroying the entire market. The anticipation of such an extreme and perverse outcome would strongly militate against limits on risk classification. More generally, the less severe inefficiencies typically associated with limits on risk classification weigh against their potential benefits.

This essay provides a unified analytical framework for systematically studying the equity-efficiency tradeoffs of legal risk-classification in competitive insurance markets. In Section 2 we develop a basic framework based on canonical adverse selection insurance markets in the Rothschild and Stiglitz (1976) and Wilson (1977) tradition. In this framework we assume information is exogenous, in the sense that individuals seeking insurance are fully informed about their risk type. While we assume that insurers are initially uninformed about risk types, we allow for the possibility that insurers can employ potentially costly tests to learn some or all of the individual’s private information. Sections 3-5 analyze the equity

³ See Joly et al (2010) for a global perspective on the regulation of genetic discrimination. See also Chiappori (2006) and Hoy and Witt (2007).

and efficiency consequences of restrictions on risk-classification in this general framework. Section 3 considers a class of environments with exogenously determined contracts in which the consequences of restrictions on risk classification are purely *distributional*. Section 4 considers a related class of environments where such restrictions have pure, and purely negative, efficiency consequences. Section 5 discusses endogenous-contract “screening” environments in which firms design contracts to induce individuals to reveal private information. Restrictions on risk classification in these environments typically involve tradeoffs between distributional and efficiency goals. Sections 6 and 7 extend our analysis to environments where risk classification is related to the endogenous choices of (potentially) insured individuals: section 6 studies environments where individuals endogenously acquire *information* about their own risk type, and section 7 studies environments where risk-classification is based on the endogenously chosen *actions* of the insured individuals. Section 8 discusses work at the frontiers of insurance market theory, and its implications for the analysis of risk classification, and suggests important directions for future research. We offer some brief conclusions in section 9.

2. The Canonical Modeling Framework

We base our analysis on the standard adverse selection model of insurance markets introduced in the economics literature by Arrow (1963) and developed by Rothschild and Stiglitz (1976) and Wilson (1977). Individuals are endowed with a lump sum of money W and face a risk of a monetary loss of size D . There is a set I of distinct types of individuals, indexed by i , who differ only in their probability p^i of experiencing the loss. For much of the analysis we focus on the case with two types, called H (igh Risk) and L (ow risk), and we assume that $p^H > p^L$ and that there is a fraction λ of H -types. More generally, we denote by $\Lambda(i)$ the population fraction of i -types

Insurance contracts are offered by risk neutral insurers. Contracts consist of a premium R , paid by the insured in both the loss and no-loss states, and an indemnity M which is paid to the insured in the loss state. An individual who purchases such a contract thus has net wealth $W - R \equiv C_1$ available for consumption if she does not experience a loss and $W - D + M - R \equiv C_2$ if she does experience a loss. We henceforth describe contracts directly in terms of the induced consumption allocations (C_1, C_2) .

Insurers are risk-neutral and, since $R = W - C_1$ and $M = D - (C_1 - C_2)$, they earn profit

$$\pi(C_1, C_2, p^i) = (W - C_1) - p^i(D - C_1 + C_2) \quad (1)$$

from selling a contract (C_1, C_2) to an i -type.

Individuals choose (at most) a single contract to maximize their expected utility,

$$V(C_1, C_2, p^i) = (1 - p^i)u(C_1) + p^iu(C_2), \quad (2)$$

where $u'(C) > 0$ and $u''(C) < 0$.

2.1 Informational environments

We consider both symmetric and asymmetric informational environments. With symmetric information, we can assume without loss of generality that firms and individuals observe the true type of each potential customer.⁴

The asymmetric information environments we focus on are those in which insurers are less informed than individuals about their risk type. (See Villeneuve, 2005, for an analysis of markets in which the insurer is more informed than the insured.) Then we can assume, without loss of generality, that individuals are fully informed about their true risk type.⁵

Completely uninformed insurers know only the population distribution $\Lambda(i)$; all individuals are indistinguishable to them. We model *partially* informed insurers as in Hoy (1982) and Crocker and Snow (1986): partially informed insurers observe an informative “signal” $\sigma \in \Sigma$ about each individual’s risk type. The signal is informative insofar as the signal-conditional probability distributions $\Lambda(i|\sigma)$ differ from $\Lambda(i)$. Signals are also referred to in the literature as “groups”, “categories”, or “classes”.

When there are two risk types, H and L , and two groups, $\sigma \in \{A, B\}$ —e.g. males and females—we take $\Lambda(H|A) = \lambda_A$ and $\Lambda(H|B) = \lambda_B > \lambda_A$ so that there is a higher fraction of low-risk types in group A ; when $\lambda_A > 0$, however, group A nevertheless contains some high-risk individuals as well; similarly, group B has relatively more high-risk types but may contain low risks as well.

⁴ A model with symmetrically known but uncertain risks p^i is isomorphic to one with symmetrically known and certain risks $\tilde{p}^i = \mathbb{E}[p^i]$.

⁵ We relax this in section 6 where we consider endogenous information acquisition.

2.2 Market outcomes

We refer to the consumption allocations $\{(C_1^i, C_2^i)\}_{i \in I}$ obtained by the various risk types in a given market environment as a “market outcome”. These are frequently (but not exclusively) referred to as “equilibrium” outcomes in the literature, and we occasionally use that terminology here as well, but the literature is somewhat inconsistent in the formal underpinnings of various equilibrium concepts (see Hellwig, 1987, for a discussion).

The market outcome that obtains in a given market environment naturally depends on the informational and the institutional features of that environment. We focus on two types of institutions, both of which are competitive and both of which involve exclusive contracting. We refer to these two types as the *fixed contracts* and *screening* cases.

In the fixed contracts case, firms compete on premium to provide an exogenously fixed indemnity M . We focus, in particular, on the full insurance case where $M = D$, although the analysis could be readily extended to a setting with $M < D$ as well. This is a useful stylized model for examining settings in which, by law or by custom, insurance contracts are standardized. We assume that the market outcomes in this case are given by the lowest-price Nash equilibrium of a game in which (a large number of) firms first set contract prices—potentially depending on the signal σ —and individuals then choose the lowest priced contract available to them.

In the screening case, firms compete on two dimensions: the price per unit of coverage R/M and on the level of coverage M . This additional flexibility allows firms to offer menus of contracts designed to differentially appeal to different risk types. A single firm might, for example, offer both a high-deductible low-premium contract and a low-deductible high-premium contract in the hopes of inducing individuals to reveal their risk type via their contract choices. Such a screening strategy is potentially useful when firms cannot observe risk type or else face regulatory restrictions on offering type-specific contracts.

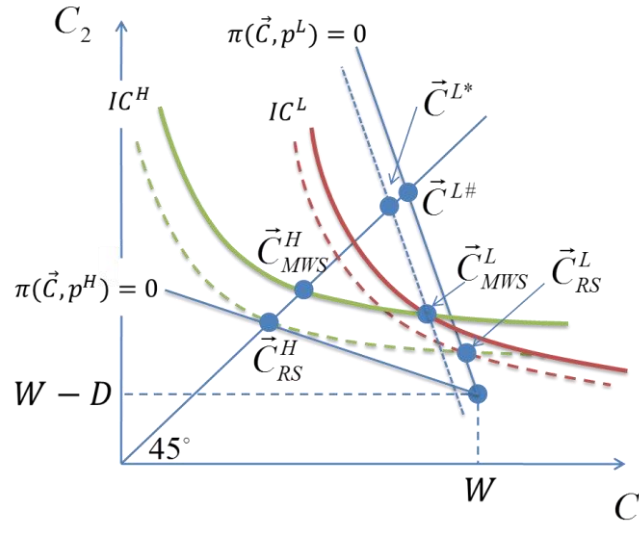


Figure 1: RS and MWS Market Outcomes

The canonical market outcome in the screening case is the Rothschild-Stiglitz (1976) equilibrium, which is depicted by the pair of points $\vec{C}_{RS}^H$ and $\vec{C}_{RS}^L$ in Figure 1 for a market with two risk types (and unobservable private information). This market outcome allocates the individually break-even, full insurance contract $\vec{C}_{RS}^H$ to the H -types and, to the L -types, the individually break-even contract $\vec{C}_{RS}^L$ that lies on the H -type's indifference curve through $\vec{C}_{RS}^H$. This indifference curve is the dashed labeled IC^H in Figure 1.⁶

The Rothschild-Stiglitz (1976) equilibrium is the only possible Nash equilibrium of a game in which multiple firms simultaneously offer individual contracts and then individuals choose their preferred contract from the set of contracts offered. Unfortunately, this allocation is *actually* a Nash equilibrium only when the fraction λ of high risk types is sufficiently high.⁷ Moreover, as Hendren (2014) demonstrates, the lack of existence of a non-trivial Rothschild-Stiglitz equilibrium is, in fact, a *generic* property of models with sufficiently rich type spaces. This makes the Rothschild-Stiglitz equilibrium concept ill-suited to general analyses of market outcomes in insurance markets.

⁶ Throughout the paper, we use IC^i to label i -type indifference curves, but it is worth noting that the underlying allocations are, in some cases, diagram-specific.

⁷ Insurers have an incentive to deviate to a pooling contract that attracts both risk types when λ is low, but no pooling contract can be a Nash equilibrium (*viz* Rothschild and Stiglitz (1976)).

The literature has employed several distinct alternative equilibrium concepts to resolve this non-existence problem, each of which predicts a unique equilibrium outcome for any set of parameters. The Riley (1979) “reactive” equilibrium coincides with the Rothschild-Stiglitz equilibrium candidate—regardless of the type distribution Λ . The Wilson (1977) E2 “foresight” equilibrium coincides with the Rothschild-Stiglitz equilibrium whenever the latter exists, and otherwise is a pooling equilibrium in which both risk types receive the L -type’s most preferred pooled zero-profit contract.

The so-called Miyazaki (1977)-Wilson (1977)-Spence (1978) (henceforth, MWS) equilibrium rests on a similar “foresight” concept, but allows subsidies across contracts.⁸ The contract pair $\vec{C}_{MWS}^H$ and $\vec{C}_{MWS}^L$ in Figure 1 qualitatively depicts a cross-subsidizing MWS market outcome in the two-type private information setting. In this market outcome, H -types get full insurance at *better* than actuarially fair rates. L -types’ contract $\vec{C}_{MWS}^L$ leaves them underinsured, but less so than with the Rothschild-Stiglitz candidate $\vec{C}_{RS}^L$. The additional insurance is valuable enough to L -types that they prefer $\vec{C}_{MWS}^L$ to $\vec{C}_{RS}^L$ in spite of its lower actuarial value—so the cross subsidies implicit in the MWS pair are Pareto improving. In the MWS market outcome, the market takes advantage of these Pareto-improving cross subsidies whenever they exist. Specifically, in the two-type model the MWS outcome is the member of the class of constrained efficient separating allocations which maximizes the expected utility of L -types subject to H -types being at least as well off as they would be with their full insurance actuarially fair contract.⁹

An alternative approach to modeling market outcomes in the insurance literature is to focus on efficient allocations (see, e.g., Crocker and Snow, 2013) rather than market equilibrium outcomes *per se*.

With symmetric information, a “first best efficient” allocation is one that maximizes the well-being $V^j(\vec{C}^j, p^j)$ of (e.g.) the j -type subject to resource feasibility,

$$\sum_{i \in I} \Lambda(i) \pi(\vec{C}^i, p^i) \geq 0, \quad (3)$$

and a set of minimum utility constraints,

⁸ See Netzer and Scheuer (2013), Mimra and Wambach (2011), and Picard (2014) for recent game theoretic foundations of the MWS equilibrium. See Dubey and Geanakoplos (2002) and Martin (2007) for discussions of the microfoundations of the Rothschild-Stiglitz-Riley equilibrium.

⁹ With I types, the MWS equilibrium can be defined recursively as the allocation which maximizes the utility of the lowest risk type subject to minimum utility constraints for higher risk types j which correspond to the utility j would achieve in an MWS equilibrium for a market consisting only of the types j and riskier; see Spence (1978).

$$V(\vec{C}^i, p^i) \geq \bar{V}^i, \quad (4)$$

for all other risk-types $i \in I$ with $i \neq j$. Every point on the first best frontier is the solution to this mathematical program for an appropriate choice of the $\bar{V}^i$ in (4).¹⁰ An allocation is on the first best efficient frontier if and only if (i) there are zero aggregate profits and (ii) each type receives full insurance (i.e., $C_1^i = C_2^i$ for all i). A social planner who observes every individual's risk type and can directly assign allocations based on this information can, in principle, implement any allocation on the first-best efficient frontier.

When individuals are privately informed about risk type and the social planner is completely uninformed about an individual's risk type, however, a “second-best” efficiency problem for the social planner is appropriate. In this second-best (or “constrained”) efficient problem, the social planner faces the additional incentive compatibility constraints:

$$V(\vec{C}^i, p^i) \geq V(\vec{C}^j, p^i) \quad \forall i, j. \quad (5)$$

These constraints reflect the fact that insofar the social planner cannot directly identify an individual's type, and thus cannot force an individual to select a contract which she finds strictly worse than some other contract which is available to another individual who is indistinguishable from them.¹¹

When the social planner is partially informed about individual risk types via a signal σ , it can condition contracts on the signal, so the incentive compatibility constraints for the constrained efficient problem are instead:

$$V(\vec{C}^{i,\sigma}, p^i) \geq V(\vec{C}^{j,\sigma}, p^i) \quad \forall i, j, \sigma \text{ such that } \Lambda(i|\sigma) > 0, \Lambda(j|\sigma) > 0. \quad (6)$$

Our approach to studying the effects of risk classification on market outcomes will be to compare market outcomes in the presence of risk classification to market outcomes in its absence. Our general, but not exclusive, focus is on the MWS outcomes.

We do not consider some other models of market outcomes discussed in the literature. For example, we do not explicitly consider models with a monopoly insurance provider (e.g., Stiglitz, 1977 and Chade and

¹⁰ This is just a definition of Pareto efficiency. As such, these minimum utility constraints should not be interpreted as participation constraints—they may be lower, e.g., than $V((W, W - D), p^i)$.

¹¹ Of course, the social planner could always choose to offer a single contract to *all* individuals, in which case constraint (5) would have no “bite.”

Schlee, 2012) or oligopolistic markets (e.g., Buzzacchi and Valletti, 2005). Nor do we consider linear pricing equilibrium, which are often used in markets where contracting is non-exclusive and individuals can buy small amounts of coverage from multiple providers simultaneously (e.g., Pauly, 1974, Hoy and Polbgorn, 2000, Villeneuve, 2003, and Rothschild, 2014).

3. Risk Classification with Purely Distributional Consequences

Policy discussions about risk classification frequently emphasize the perceived distributional benefits of restricting firms from employing risk classification and downplay or ignore the potential efficiency costs of such restrictions. We first consider a (rather restrictive) setting in which risk-classification does not have any efficiency consequences and this emphasis is appropriate.

This setting is characterized by the following assumptions:

- (i) there are **two risk types**, H and L , with $p^H > p^L$, and a fraction λ of H -types;
- (ii) there is **symmetric information**: insurers and individuals can both observe type;
- (iii) insurance contracts are **fixed full insurance contracts**—so firms compete only on price;
- (iv) there is a **mandatory purchase** requirement: each individual must buy exactly one contract.

Figure 2 depicts the Nash equilibrium market outcomes.¹² When insurance providers classify based on observable risk type, each type $i \in \{H, L\}$ pays its type-fair premium $R^i = p^i D$ in exchange for full indemnification $M = D$ of the loss. The associated consumption allocations are labeled $\vec{C}^{i\#}$, $i = H, L$. When insurance providers do not employ risk classification based on risk type, both types of individuals pay the pooled-fair premium $\bar{R} = (\lambda p^H + (1 - \lambda)p^L)D \equiv \bar{p}D$, which yields consumption $\vec{C}^{P\#}$.

These are the unique Nash equilibria of the pricing games that arise, respectively, with legal and with banned risk classification. In particular, risk classification will be used by the market if it is permitted since (e.g.), at the no classification equilibrium $\vec{C}^{P\#}$, a deviating firm could make positive profits offering a slightly less expensive contract (just up and to the right from $\vec{C}^{P\#}$) which is available only to L -types. This will attract all L -types, leaving only the H -types buying $\vec{C}^{P\#}$ and rendering that contract unprofitable. As such, the no-classification market outcome is likely to arise only when there are explicit

¹² Note that we can use Nash equilibrium in this game, since firms compete only on price. When firms also compete on product design, we will focus on the MWS equilibrium concept.

legal restrictions on the use of risk classification—and only then in the absence of other regulatory interventions (e.g., risk-adjustment schemes) that mitigate firms’ private incentives to “cream skim” the low risk types.

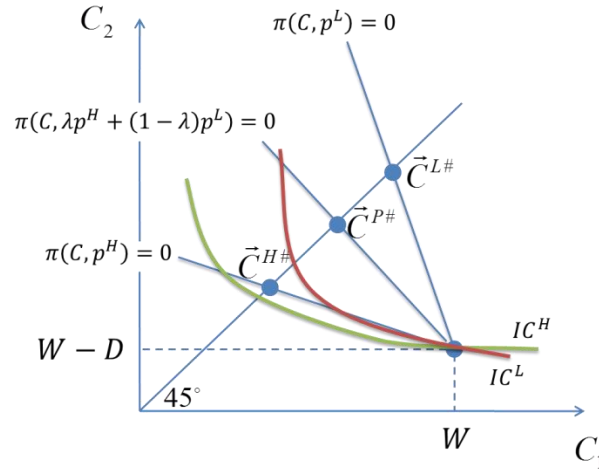


Figure 2: Purely Redistributive Risk Classification

Since the market outcomes both with and without risk classification are first-best efficient, banning risk classification has *purely* distributional consequences in this setting: it improves the well-being of the relatively worse-off *H*-types at the expense of the relatively well-off *L*-types without causing any inefficiency. With typical distributional preferences, bans on risk classification are therefore welfare improving in this institutional setting.

To elaborate on this point, it is useful to consider an alternative interpretation of Figure 2 and the preceding argument. Suppose that risk classification is permitted, but individuals are all identical and are originally *uninformed* about their risk types; consequently, all buy insurance at the pooled fair price and consume at the point $\vec{C}^{P\#}$. Now suppose a new test is developed that will perfectly reveal their type—either *H* or *L*. If individuals who take the test will be offered insurance at their type-fair price, then individuals, being risk averse, will choose *not* to take the test: taking the test has no effect on their (full) coverage, but introduces “premium risk” that reduces their expected utility insofar as they are risk averse. Requiring everyone to take such a test would be efficiency reducing.¹³

¹³ This omits the possibility that tests convey treatment-relevant information—in which case undertesting due to premium risk can cause additional inefficiencies. Section 6 discusses endogenous informational environments in more detail.

Turning this around, if types are instead *known*, at the time of insurance contracting, then type *is itself a risk* from an ex-ante point of view before types are realized. From this ex-ante point of view, then, the premium risk associated with the realization of type is itself inefficient. So a ban on risk classification that has the effect of eliminating premium risk is efficiency enhancing from an *ex-ante* point of view.

The assumptions (i)-(iv) underlying the analysis in this section are obviously quite restrictive, but assumptions (i) and (ii) are easily relaxed: the same basic result would hold if there were more risk types, or if insurers could observe an imperfect signal of individuals' privately known risk type. On the other hand, as we discuss in the following section, if contracts are not regulated to provide fixed, full coverage per (iii) and there is no individual purchase requirement, there will typically be negative efficiency consequences of imposing legal restrictions on risk classification in screening settings.

The mandatory purchase assumption (iv) may or may not be essential, depending on specifics of individual preferences and on the distribution of risk types. In the situation depicted in Figure 2, (iv) is *not* essential since the pooled-fair full insurance point $\vec{C}^{P\#}$ lies strictly above the "autarkic" indifference curve IC^L that passes through the no-insurance consumption point $(W, W - D)$. This means that, if risk classification were banned, both individuals would *voluntarily* choose $\vec{C}^{P\#}$ rather than forego insurance. Assumption (iv) is critical, on the other hand, if L -types prefer foregoing insurance to $\vec{C}^{P\#}$.

4. Risk Classification with Pure Efficiency Consequences

The assumptions in Section 3 restrict the behavior of firms and individuals sufficiently to ensure that the only effect a risk-classification ban will have is a change in the prices paid by different individuals. Indeed, under these assumptions, a ban on risk classification will result in universal insurance at a pooled fair price. A ban on risk classification *in the presence of other strict* regulations, namely a purchase mandate and strong minimum coverage requirement is, therefore, not only unproblematic, but, in fact, yields an outcome—universal full insurance against risk *and* full insurance against classification risk—that is particularly normatively compelling.

This section shows that the welfare effects of risk classification bans can be quite different in the absence of other strict regulations.

When the mandated purchase assumption (iv) is relaxed, for example, a ban on risk-classification can reduce efficiency without any beneficial distributional consequences at all. To wit: when the fraction λ of high-risk types is sufficiently high the pooled fair full insurance contract $\vec{C}^{P\#}$ lies close to $\vec{C}^{H\#}$ and $V(\vec{C}^{P\#}, p^L) < V(W, W - D, p^L)$; under a ban, then, any premium at which L -types will buy insurance will be strictly unprofitable when it is also sold to H -types. Bans therefore lead L -types to exit the insurance market and make them strictly worse off than they would be in the absence of a ban, without improving the lot of the H -types, who end the exact same contract $\vec{C}^{H\#}$ as they obtain in a risk-classifying regime.¹⁴

The negative efficiency consequences of a ban on risk-classification in a fixed, full insurance contract, voluntary-purchase environment can be quite dire. Suppose, for example, that there is a continuum of risk types with risk probabilities p^i uniformly distributed on $[0,1]$. For any given premium $R < D$, only individuals with

$$p^i \geq q(R) = \frac{u(W) - u(W - R)}{u(W) - u(W - D)} \quad (7)$$

will voluntarily purchase insurance. The break-even premium for this buyer pool is $D(1 + q(R))/2$. When $D(1 + q(R))/2 > R$ for all $R < D$ (e.g., with $u(x) = x - kx^2$ and $2Wk < 1$, which implies $2xk < 1$ and hence $u'(C) > 0$), then there is no premium $R < D$ at which firms are willing to sell contracts. The market completely unravels via an “adverse selection death spiral”¹⁵ and no insurance at all is provided.

As we discuss in the next section, without fixed full insurance contracts (i.e., when firms can freely choose coverage levels), firms may be able to use screening mechanisms to prevent a complete collapse

¹⁴ The Affordable Care Act contains a purchase mandate precisely to prevent this sort of exit-based unraveling, although Feldstein (2013) argues that the fine for violating the mandate is insufficient to prevent exit. See also Handel et al (2013), who use a quantitative equilibrium model to predict both the degree of unraveling through non-compliance with the mandate, per Feldstein, and to predict the degree of “unravelling” to pooling lower-coverage “bronze” plans.

¹⁵ See, e.g., Akerlof (1970), Cutler and Reber (1998) and Strohmeier and Wambach (2000).

of the market for insurance,¹⁶ but even in this case, risk classification bans can still be purely efficiency reducing.¹⁷

5. Risk Classification with Efficiency/Redistribution Trade-Offs

Risk classification in screening environments will frequently involve a non-trivial tradeoff between the efficiency and distributional equity consequences. In this section, we use the MWS market outcomes to illustrate this tradeoff.¹⁸ We then discuss various approaches to the welfare analysis of risk classification bans in the presence of these tradeoffs.

5.1 Efficiency and distributional effects with MWS market outcomes

Figure 1 can be used to analyze the two-type, symmetric information case. With risk classification, individuals would receive full insurance at their actuarially fair full insurance premium, resulting in the allocation $\vec{C}^{L\#}$ for L -types and $\vec{C}_{RS}^H$ for H -types. When risk-classification is banned, the market implements the MWS contracts $\vec{C}_{MWS}^i$. Since, as drawn, this market outcome involves cross-subsidies from the L -types to the H -types, H -types are *strictly* better off with banned risk-classification while L -types are strictly worse off (but better off than with the Rothschild-Stiglitz candidate $\vec{C}_{RS}^L$.) For a social planner who would like to redistribute from L -types to H -types, banning risk classification thus has desirable distributional consequences.

The ban, however, also has unambiguously *negative* efficiency consequences: it moves the economy from a first-best efficient allocation to one which is not first best efficient. Specifically, it is inefficient because replacing the L -type's allocation $\vec{C}_{MWS}^L$ with the full insurance $\vec{C}^{L*}$ depicted in Figure 1 is information- and resource-feasible and strictly improves the welfare of L -types. Moreover, by imposing a properly calibrated tax/subsidy on contracts sold to L/H types, a social planner could induce a decentralized market to implement the allocation $(\vec{C}_{MWS}^H, \vec{C}^{L*})$. So, while a social planner could, on

¹⁶ Hendren (2014, Theorem 1) provides a precise characterization of the condition under which screening mechanisms (broadly interpreted) can and cannot prevent a death spiral.

¹⁷ Buchmueller and DiNardo (2002) look empirically at the consequences of community rating (a ban on risk classification) in the small group and individual health insurance market in New York State. They provide evidence consistent with this sort of an effect: they find no evidence of a significant reduction in the fraction of individuals with insurance, but they do find evidence of a significant shift towards less generous insurance coverage.

¹⁸ See Hoy (2006) for a related analysis using the Wilson (E2) equilibrium concept.

distributional grounds, reasonably prefer a straight ban on risk classification to the pure free market outcome with legal risk classification, any welfarist social planner would strictly prefer a regime with legal risk classification and such a tax/subsidy scheme.

Crocker and Snow (1986) show that the same conclusions apply when risk classification is only partially informative. They consider the setting with two risk types and two partially informative signals A and B , such that $\Lambda(H|B) = \lambda_B > \lambda = \Lambda(H) > \lambda_A = \Lambda(H|A)$, with $\lambda_B < 1$ and $\lambda_A > 0$ (i.e., signal B indicates that an individual is assigned to the category with a higher fraction of high-risk types, and signal A indicates the category with more low-risk types). To illustrate their (more general) result, consider the MWS market outcomes. Without risk classification, the outcome $(\vec{C}_{MWS}^H, \vec{C}_{MWS}^L)$ is a standard two-type MWS equilibrium for an economy with a fraction λ of high-risks. With risk classification, the market outcome provides individuals in group σ ($\sigma \in \{A, B\}$) with the consumptions $(\vec{C}_{MWS}^{H,\sigma}, \vec{C}_{MWS}^{L,\sigma})$ associated with the MWS outcome for a two risk-type economy with the *group-specific* fractions λ_σ of high-risks. There are three possibilities:

- (1) The A -group market outcomes $(\vec{C}^{H,A}, \vec{C}^{L,A})$ do not involve cross subsidies from low to high risks.
- (2) The A -group market outcomes $(\vec{C}^{H,A}, \vec{C}^{L,A})$ involve cross subsidies and the no risk-classification market outcomes $(\vec{C}^H, \vec{C}^L)$ do not involve cross subsidies.
- (3) The no risk-classification market outcomes $(\vec{C}^{H*}, \vec{C}^{L*})$ involve cross subsidies.

These three possibilities are exhaustive because (fixing preferences and other parameters in the standard two-type setting) there exists a cutoff $\hat{\lambda} \in (0,1)$ such that the MWS market outcomes involve cross subsidies if and only if $\lambda < \hat{\lambda}$ (viz Crocker and Snow 1985a). We use this fact extensively below. We also use the fact that the well-being of *each* type is decreasing in λ for $\lambda \leq \hat{\lambda}$. Intuitively: a higher fraction of L -types means more scope for Pareto-improving cross subsidies.

In case (1), $\hat{\lambda} \leq \lambda_A < \lambda < \lambda_B$, and market outcomes coincide with the Rothschild-Stiglitz equilibrium consumptions both with and without risk classification. In this case, banning risk classification has neither efficiency nor distributional consequences—indeed, it has no economic effects at all.

In case (2), $\lambda_A < \hat{\lambda} \leq \lambda < \lambda_B$ and the market outcome with risk classification Pareto dominates the market outcome without risk classification: Category B individuals receive the same Rothschild-Stiglitz separating contracts with and without risk classification, while the Category A individuals take

advantage of Pareto-improving (within category) cross-subsidies and are strictly better off with risk classification. In this case, risk classification has only (beneficial) efficiency consequences.

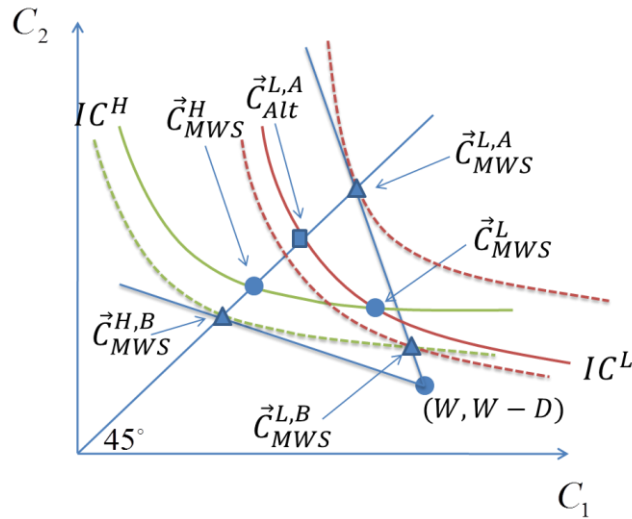


Figure 3: Risk Classification Can Have Both Efficiency and Distributional Consequences

In case (3), $\lambda_A < \lambda < \hat{\lambda}$. Figure 3 illustrates the special case with $\lambda_A = 0$ and $\lambda_B \geq \hat{\lambda}$. Triangles denote the category- σ and type- i contracts $\vec{C}_{MWS}^{i,\sigma}$ that obtain when risk classification is legal, and circles denote the type- i contracts $\vec{C}_{MWS}^i$ that obtain when classification is banned. As shown in the figure, risk classification has distributional consequences. Relative to a ban, risk classification makes category- A individuals (all of whom are L -types since $\lambda_A = 0$ here) strictly better off: it moves them to the higher dashed indifference curve. It also makes both types of category- B individuals worse off, moving them onto the lower dashed indifference curves.

Risk classification bans *also* have efficiency consequences in this case. To see this (again for the special case illustrated in Figure 3) note that category is observable and that category A is perfectly indicative of type L . So, starting from the banned-classification market outcome, it is informationally feasible, and less costly (profit increasing), to move the A category individuals from $\vec{C}_{MWS}^L$ to the allocation $\vec{C}_{Alt}^{L,A}$ depicted in Figure 3. Hence, the no risk classification outcome is not second-best efficient. The second best inefficiency of the no risk classification outcome extends to the more general case where $\lambda_A > 0$ (and arbitrary $\lambda_B > \lambda$).¹⁹

¹⁹ To wit: the no classification outcome involves cross subsidies (as $\lambda < \hat{\lambda}$). The magnitude of these cross subsidies is effectively chosen to balance the subsidy's incentive constraint-easing benefits to L -types with the cost, to L -types, of transferring resources to H -types. Lowering the fraction of H -types lowers the resource cost of any given

5.2 Welfare analysis with distributional and efficiency effects

There are three broad approaches to welfare analysis when there are both efficiency and distributional consequences (as in case (3) above). One approach is to adopt an explicit social welfare function or family of social welfare functions which assign welfare weights to the different types. This approach can be viewed as imposing a well-specified trade-off between the distributional benefits and the efficiency costs. Hoy (2006), for example, uses an explicit utilitarian social welfare function (and his results would obviously extend to a broad class of social welfare functions), and Einav and Finkelstein (2011) use a deadweight-loss based formulation that is equivalent to a utilitarian social welfare function with the particular cardinalization of individual utilities that weighs willingness-to-pay equally across individuals.²⁰

A second approach is to explicitly and separately quantify the efficiency and distributional consequences; this approach is useful when different policy analysts differ in their perception about the optimal trade-off between distribution and efficiency concerns. See, e.g., Finkelstein et al. (2009) who apply this approach to the analysis of gender classification in the U.K. compulsory annuity market).

A third approach is to argue that policymakers should focus exclusively on the efficiency consequences. Of course, in the simple models considered in this paper, one way to achieve full efficiency is to implement mandatory full insurance at a pooled price for all individuals—either through strict regulations or via direct government provision. But for a policymaker who is committed to market solutions, there are settings in which restricting attention to the efficiency consequences of risk classification is fruitful, even if that policymaker cares about equity as well as efficiency. Consider, for example, the constraint inefficient market outcome $(\tilde{C}_{MWS}^H, \tilde{C}_{MWS}^L)$ which obtains under a ban on risk classification, as depicted in Figure 3. Under certain distributional preferences, this market outcome could be preferred to the market outcome that obtains with legalized risk classification. Nevertheless,

amount of incentive constraint easing, and thus tilts the balance towards additional cross subsidies. This means that at $(\tilde{C}_{MWS}^H, \tilde{C}_{MWS}^L)$ in Figure 3, additional cross subsidies would be Pareto improving within category *A* even when $\lambda_A > 0$. There is therefore an alternative allocation that could be offered in an incentive-compatible way to the category *A* individuals that would increase the well-being of both category-*A* types at the same net actuarial cost for the category. It is therefore information- and resource-feasible to make the category *A* individuals strictly better off without changing the allocations to category-*B* individuals.

²⁰ See also Thomas (2008), who invokes a “loss coverage maximization” criterion. In contrast to the tradeoffs between the ex-ante expected *utility* levels of different individuals that is invoked by a weighted-average-of-expected-utility social welfare function, loss coverage maximization invokes an explicit tradeoff between ex-ante *coverage* levels across high and low risk types. Specifically, the criterion assigns equal weight to equal expected losses.

the fact that $(\vec{C}_{MWS}^H, \vec{C}_{MWS}^L)$ is (constrained) inefficient *means* that there is some Pareto improving alternative that a social planner facing the same informational constraints as the market could implement. Crocker and Snow (1986) (building on their earlier work (1985b)) show, moreover, that these Pareto improvements can be decentralized via a set of contract (and category) specific taxes, at least when it is costless for firms to observe the categorical signals σ .

Rothschild (2011) extends Crocker and Snow's (1986) argument by showing that the provision of partial social insurance is an *alternative* way to decentralize Pareto improvements associated with a shift from a banned classification to a legalized classification regime. His argument boils down to a simple observation: removing bans on risk classification is potentially problematic only insofar as it "undoes" the cross-subsidies provided by the market to individuals in the higher risk category. Coupling a removal of ban with another policy to "lock in" this cross subsidy moots these concerns. Rothschild (2011) shows that a simple-to-implement partial social insurance scheme can accomplish this "lock in".

Figure 4 illustrates the argument. The pair $(\vec{C}_{MWS}^H, \vec{C}_{MWS}^L)$ depicts the MWS outcome with banned categorization; as illustrated, the market involves cross subsidies from L -risks to H -risks, so the market outcome falls under case (3), and, as described above, risk classification has both distributional and efficiency consequences.

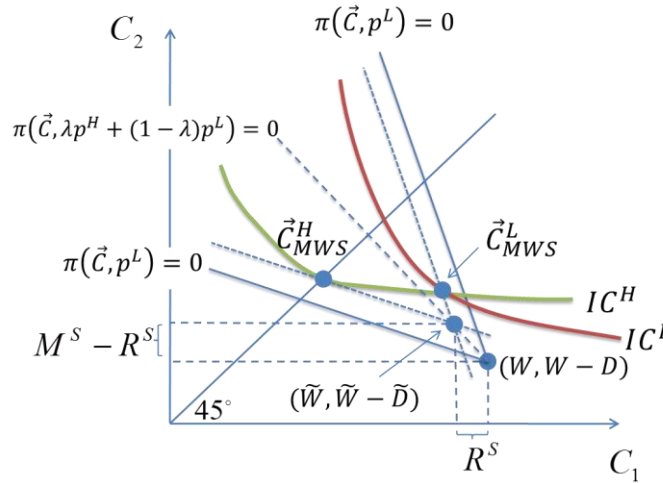


Figure 4: Partial Social Insurance is Preferable to a Ban on Risk Classification

Now suppose the government implements a social insurance policy providing the indemnity of M^S and the (mandatory) premium R^S depicted in Figure 4. This policy is the pooled-fair partial social insurance policy that yields the same L -type to H -type resource transfer as the contract pair $(\vec{C}_{MWS}^H, \vec{C}_{MWS}^L)$. With

this social insurance in place, individuals have an effective wealth of $\tilde{W} = W - R^S$ and face an effective loss of $\tilde{D} = D - M^S$, and there is still scope for supplemental private insurance. In fact, the provision of this partial social insurance does not affect the allocations $\tilde{C}_{MWS}^i$ provided under a risk classification ban at all: the MWS market outcome with endowment wealth $\tilde{W}$ and loss $\tilde{D}$ remains $(\tilde{C}_{MWS}^H, \tilde{C}_{MWS}^L)$. But now these same outcomes are implemented with zero *private market* cross subsidies from *L*-risks to *H*-risks. Indeed, with endowment wealth $\tilde{W}$ and loss $\tilde{D}$, the population fraction λ of high risks is exactly the “cutoff” fraction $\tilde{\lambda}$ between an MWS equilibrium with and without no cross subsidies. Hence, with the social insurance policy in place, the MWS equilibrium falls into case (2) and, as described above, introducing risk classification is purely efficiency improving in case (2).²¹

In an otherwise unregulated setting, Crocker and Snow (1986) and Rothschild’s (2011) arguments imply that even when banning risk classification furthers some distributional objective, imposing a ban is a sub-optimal way to achieve this objective.

As such, any welfarist social planner will prefer the outcomes in *some* feasible regime with legal risk classification to the banned risk-classification market outcomes. Insofar as these sorts of taxes are within the purview of the social planner, then, one can conclude that bans on risk-classification are never strictly desirable, regardless of distributional concerns.

6. Endogenous Informational Environments

In the preceding analysis, we modeled individuals as having a fixed amount of information about their own risk type (specifically, full information). In many applications, the assumption of static information is unrealistic, as individuals may, for example, learn more information about their own risk types through screening tests or medical exams. With the advent and development of genetic screening,

²¹ The argument is general. It does not rely on the MWS equilibrium concept; it applies even when risk-classification is costly, and it does not require the social planner to know anything about the costliness or informativeness of the risk-classification technology. It only requires that the social planner knows the magnitude of the cross subsidies that obtain in the absence of risk classification. Of course, banning risk classification could still be preferred, on distributional grounds, to a status quo policy of legal risk classification. Rothschild’s (2011) argument shows that there is some social insurance policy which is Pareto preferred to such a ban, but the social planner may not know enough *under the status-quo regime* to compute this Pareto preferred policy. A lack of necessary information is particularly likely when the risk classification technology is rapidly changing, as with genetic testing; as such, this argument—and indeed, the model underlying it—should be used with extreme caution.

analyses of insurance purchases in these sorts of “dynamic” informational environments are increasingly important. We show in this section that the basic conclusion that risk classification bans are typically associated with negative *static* efficiency consequences extends to standard models of these environments. The conclusion that risk-classification bans may involve tradeoffs between distributional equity and efficiency, when unaccompanied by other policy interventions, requires some nuanced reinterpretation, however. To see why, consider a simple setting in which individuals are *initially* uninformed but later learn about their risk type and then purchase insurance. From an *interim* point of view *after* individuals have learned their type but *before* they have purchased insurance, a ban on risk classification may (beneficially) redistribute from low to high risk types. But from the *ex-ante* point of view before learning their type, this “redistribution” provides insurance against classification risk and can be regarded as (ex-ante) efficiency enhancing. So, in dynamic settings, it may make more sense to think of risk-classification as imposing a tradeoff between *ex-ante* and *interim* efficiency.²²

Dynamic informational environments are more challenging to analyze because of the interaction between the availability of tests, testing decisions, and insurance opportunities. In particular, an individual’s *decision* to learn more about her risk type through a test will likely depend on how that information will affect her insurance options, and aggregated individual decisions will in turn affect the insurance options available.

Dionne et al. (2013) provide a comprehensive analysis of many of these issues (see also Crocker and Snow (1992) and Doherty and Thistle (1996)). We build on their basic analysis here, but we modify it in two ways. First, we focus specifically on the implications of dynamic information for *risk classification*. Second, we employ the MWS market outcome concept instead of the Rothschild-Stiglitz concept, thereby resolving some non-existence problems that can arise in their analyses.

The basic framework is a simple extension of the two-type framework discussed above. There are two risk types H and L and two possible informational states for individuals, “informed” and “uninformed”. Informed individuals know their risk type; uninformed individuals (U -types) do not. We let κ denote the fraction of individuals who are informed; κ is an endogenous quantity insofar as individuals can choose whether or not to become informed. We denote by $\kappa_0 \in [0,1)$ the exogenous initial fraction of

²² This is not always a precise distinction. For example, interim inefficiency means that there is a potential intervention that effects a Pareto improvement at the interim stage. Any such improvement would, of course, be desirable at the *ex-ante* stage as well. As such, it is not so much a true “tradeoff” between interim and ex-ante efficiency as it is a question of the implementation of the potential Pareto improvement at the interim stage.

informed individuals. The fraction of informed individuals who are H -types, denoted again by λ , is assumed independent of κ . Finally, there is a fixed utility cost $\tau \geq 0$ for an uninformed type to take a test and become informed. U -types are uncertain about their true risk type and, if they remain uninformed, have the same preferences as an individual with risk type $p^U = \lambda p^H + (1 - \lambda)p^L$.

We first describe the MWS market outcomes in four cases, which differ in the availability of information to insurers (or to a notional social planner). In each of these cases, an MWS market outcome consists of both the contracts $\vec{C}_{MWS}^i$ received by $i = H, L, U$ types *and* the endogenous fraction κ of informed types. We then use these to analyze the welfare consequences of banning risk classification.

6.1 Symmetric information

If information is completely symmetric, the market outcomes are simple: each type i ($i = H, L, U$) receives a type-specific actuarially fair full insurance contract $\vec{C}^{i\#}$ (viz Figure 5). Since becoming informed does not change the expected (actuarial) *value* of U -types' consumption but *does* introduce welfare reducing uncertainty—in the price of the contract received— U -types have no incentive to become informed, and $\kappa = \kappa_0$.

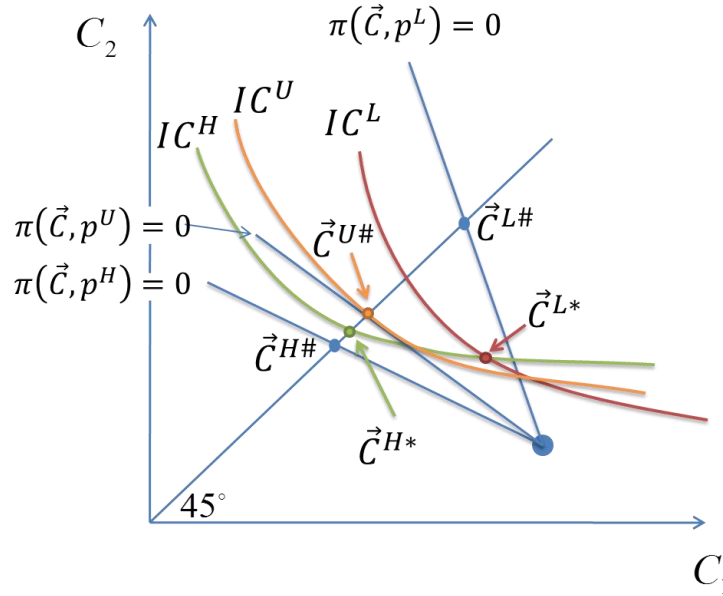


Figure 5: Market Outcomes with Private Test Results and Public Information Status

6.2 Private test results, public information status

Figure 5 depicts the market outcome if insurance firms can observe whether an individual has been tested or not—and can condition their contracts on this observation—but do not observe the outcome of that test. In this case, U -types again get their full-insurance actuarially fair contract $\vec{C}^{U\#}$. H - and L -types get the contracts $\vec{C}^{H*}$ and $\vec{C}^{L*}$ associated with the MWS market outcome for a two risk-type economy. (As drawn, these involve cross subsidies; our analysis does not depend on this.)

As in the symmetric information case, uninformed individuals again do *not* have an incentive to become informed in this market outcome: taking the test and becoming informed does not affect the expected *value* of the consumption provided by insurance, but it introduces *two* types of risk: classification risk, as in the symmetric case, and the uninsured accident risk associated with the fact that $\vec{C}^{L*}$ provides less than full insurance. Hence, $\kappa = \kappa_0$.²³

6.3 Private information status, verifiable test outcomes

If type and informational status are private information, but test outcomes can be verifiably reported for individuals who have taken the test, then L -types will reveal their type and receive full fair insurance. H -types will not reveal their test status and will be indistinguishable from U types.²⁴

Fixing any $\kappa \in [\kappa_0, 1)$, the market for the H and U types will be a standard two-type MWS outcome, with U -types playing the role usually played by L -types. We denote by $\vec{C}^{i\$}(\kappa)$ the contract received by i -types ($i = U, H$). Since the ratio of H - to U -types is $(\lambda\kappa)/(1 - \kappa)$, a higher κ reduces the scope for utility-improving cross subsidies from the U - to the H -types. There is some $\bar{\kappa}$ such that there are zero cross subsidies when $\kappa \geq \bar{\kappa}$ and positive cross subsidies when $\kappa < \bar{\kappa}$. Hence, when $\kappa < \bar{\kappa}$, $V(\vec{C}^{U\$}(\kappa), p^U)$ is strictly decreasing in κ and when $\kappa \geq \bar{\kappa}$, $V(\vec{C}^{U\$}(\kappa), p^U)$ is independent of κ .

The value of information to U -types is:

²³ If information status is public and test results are private but verifiable, L -types will then have an incentive to verifiably reveal their test results and receive their full-insurance actuarially fair contract $\vec{C}^{L\#}$. Moreover, firms will then infer that informed types who did not reveal their risk types are H -types; they will therefore receive *their* full-insurance actuarially fair contract $\vec{C}^{H\#}$. This case thus effectively reduces to the symmetric information case.

²⁴ This informational regime would naturally apply where testing is, by law or custom, anonymous. Then individuals can credibly reveal test results but *cannot* credibly “reveal” that they have not been tested.

$$\begin{aligned}
I(\kappa) &= \lambda[V(\vec{C}^{H\$}(\kappa), p^H) - V(\vec{C}^{U\$}(\kappa), p^H)] + (1 - \lambda)[V(\vec{C}^{L\#}, p^L) - V(\vec{C}^{U\$}(\kappa), p^L)] \\
&= (1 - \lambda)[V(\vec{C}^{L\#}, p^L) - V(\vec{C}^{U\$}(\kappa), p^L)],
\end{aligned} \tag{8}$$

where the first equality uses the fact that individuals will not reveal the test results if they turn out to be H , and where the second equality follows from the binding incentive compatibility constraint in an MWS equilibrium.

Figure 6 qualitatively plots $I(\kappa)$ in the non-trivial $\kappa_0 < \bar{\kappa}$ case. It is continuous on $\kappa \in [\kappa_0, 1)$ but jumps up discretely at $\kappa = 1$ since the market will offer no contract to (non-existent) U -types there, and choosing to remain uninformed would force U -types into contract $\vec{C}^{H\#}$.

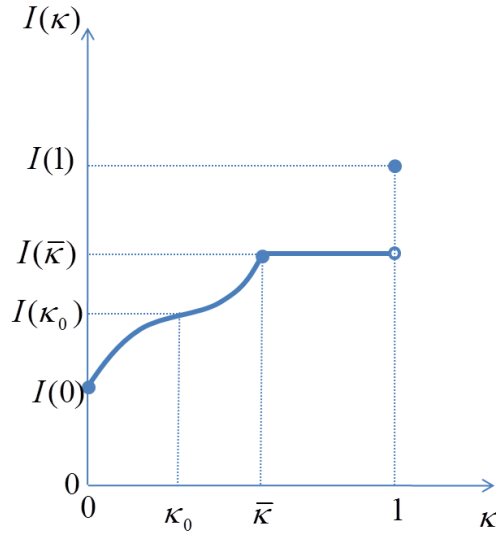


Figure 6: The Value of Information with Unobservable but Verifiable Tests

The equilibrium values of κ can be found for various testing costs τ . When $\tau < I(\kappa_0) < I(1)$, all individuals want to get tested no matter what other individuals do. The market outcome therefore involves $\kappa = 1$ and involves all individuals getting tested and receiving their type-specific actuarially fair full insurance contracts. When $I(1) < \tau$, nobody wants to get tested and the unique equilibrium has $\kappa = \kappa_0$. For intermediate values of τ there are multiple equilibrium values of κ . This multiplicity arises from a coordination problem: as Figure 6 illustrates, the incentive to become informed increases as more and more individuals become informed. When $I(\bar{\kappa}) < \tau \leq I(1)$, there are two equilibria: either all individuals coordinate on remaining uninformed (so that $\kappa = \kappa_0$) or they coordinate on becoming informed and $\kappa = 1$. When $I(\kappa_0) \leq \tau < I(\bar{\kappa})$, there is, additionally, a third (unstable) equilibrium κ : the

unique intermediate value satisfying $I(\kappa) = \tau$. Whenever there are multiple equilibria, they can be Pareto ranked, with lower κ equilibria dominating higher κ equilibria.

6.4 Purely private information

When information is purely private, the market outcome can be found by first considering the market outcome for any *given* $\kappa \geq \kappa_0$ and then looking for equilibrium values of κ .

When $\kappa \in (0,1)$, there are effectively three unobservably distinct types. Figure 7 provides a qualitative depiction of the MWS market outcome $(\vec{C}^{L'}(\kappa), \vec{C}^{U'}(\kappa), \vec{C}^{H'}(\kappa))$. In this outcome, H -types receive full insurance, and there are two binding incentive compatibility constraints:

$$V(\vec{C}^{H'}(\kappa), p^H) = V(\vec{C}^{U'}(\kappa), p^H) \text{ and } V(\vec{C}^{U'}(\kappa), p^U) = V(\vec{C}^{L'}(\kappa), p^U). \quad (9)$$

Figure 7 depicts the case where L -types provide cross subsidies to the U -types but H -types don't receive a cross subsidy; other configurations are possible.

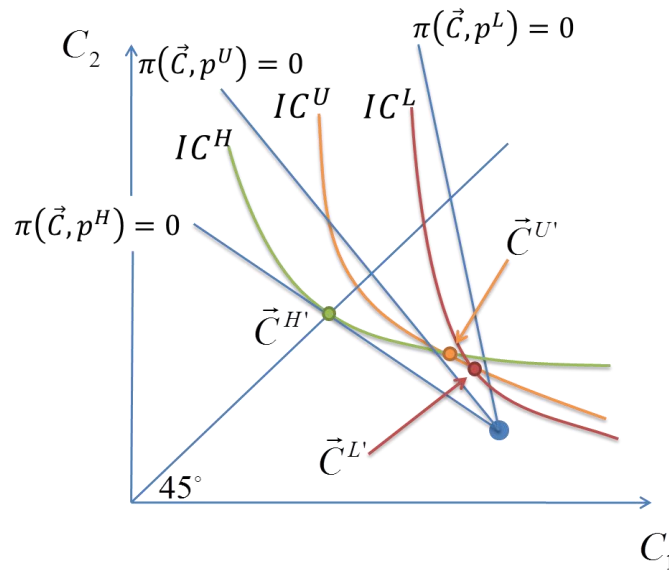


Figure 7: A 3-type MWS Equilibrium

For any $\kappa \in (0,1)$, the value of information $I(\kappa)$ satisfies

$$\begin{aligned} I(\kappa) &\equiv (1 - \lambda)[V(\vec{C}^{L'}(\kappa), p^L) - V(\vec{C}^{U'}(\kappa), p^L)] + \lambda[V(\vec{C}^{H'}(\kappa), p^H) - V(\vec{C}^{U'}(\kappa), p^H)] \\ &= (1 - \lambda)[V(\vec{C}^{L'}(\kappa), p^L) - V(\vec{C}^{U'}(\kappa), p^L)] \geq 0, \end{aligned} \quad (10)$$

with strict inequality if and only if $\vec{C}^{L'} \neq \vec{C}^{U'}$. The equality and inequality on the second line respectively follow from the binding H -type incentive constraint and the combination of the U -type incentive constraint and the single crossing property (L -types' indifference curves are steeper than U -types').

We now establish that an equilibrium always exists. If $I(\kappa_0) \leq \tau$, then there is obviously an equilibrium with $\kappa = \kappa_0$. If $I(\kappa_0) > \tau$, the following three facts hold: (i) $I(\kappa)$ is continuous, (ii) $\lim_{\kappa \rightarrow 0} I(\kappa) > 0$, and (iii) $I(\kappa) = 0$ for all $\kappa \in [\bar{\kappa}, 1)$ for some $\bar{\kappa} < 1$.²⁵ Hence, there is a $\kappa \in (\kappa_0, 1)$ such that $I(\kappa) = \tau$ which, together with the corresponding allocation $(\vec{C}^{L'}(\kappa), \vec{C}^{U'}(\kappa), \vec{C}^{H'}(\kappa))$ constitutes an equilibrium.²⁶

6.5 The inefficiency of risk classification bans with endogenous information

Risk classification is only relevant when firms potentially have access to some information which they can use for classification—i.e., in the informational environments of sections 6.1, 6.2, and 6.3. Banning risk classification in each of these settings effectively makes information purely private, as in section 6.4. We now show that such a ban leads to economic inefficiencies in each of the three settings. To do so, we explicitly construct incentive and information compatible contract menus that could, in principle, be offered by a social planner with no more information than insurance firms and that would Pareto improve on the market outcome that obtains without risk classification.

²⁵ Fact (i) follows from the continuity in κ of the MWS allocations. Fact (ii) can be proved using continuity and the observation that the solution to the 3-type MWS program (per Spence 1978) when $\Lambda(L) = \Lambda(H) = 0$ has $\vec{C}^{L'} \neq \vec{C}^{U'}$ (the allocation corresponds with the Rothschild-Stiglitz allocation for these two types). Fact (iii) can be proved by observing (1) that the MWS outcome in 3-type problems with a *zero* measure of the intermediate U -types involves pooling U -types and L -types; (2) that separating U -types from L -types in this allocation *strictly* lowers the L -types utility in this solution; and (3) by continuity, the same is true when the measure of U -types is sufficiently close to 0.

²⁶ Note that this is in contrast to Doherty and Thistle (1996), who use the Rothschild-Stiglitz equilibrium. The key difference being that the MWS equilibrium is continuous in κ at $\kappa = 1$ whereas the utility of the L -type jumps discretely at $\kappa = 1$ in the RS equilibrium.

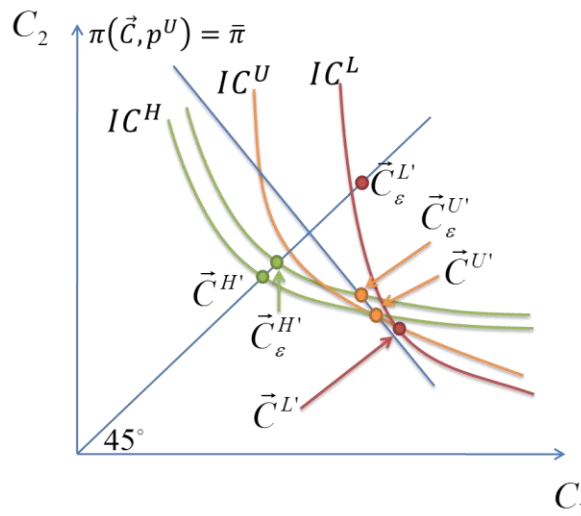


Figure 8: The Inefficiency of Banning Risk Classification with Endogenous Information

Suppose first that when risk classification is banned, the market outcome involves $\kappa < 1$, so there are some uninformed individuals. A possible market outcome $(\vec{C}^{L'}, \vec{C}^{U'}, \vec{C}^{H'})$ is depicted in Figure 8 (which remains deliberately agnostic about the cross-subsidies.) Figure 8 also depicts an alternative menu of contracts $(\vec{C}_{\varepsilon}^{L'}, \vec{C}_{\varepsilon}^{U'}, \vec{C}_{\varepsilon}^{H'})$. This menu is constructed in three sequential steps:

1. $C_{\varepsilon,1}^{H'} = C_{\varepsilon,2}^{H'} = C_1^{H'} + \varepsilon = C_2^{H'} + \varepsilon$, so H -types are offered slightly more generous full insurance coverage.
2. $V(\vec{C}_{\varepsilon}^{U'}, p^H) = V(\vec{C}_{\varepsilon}^{H'}, p^H)$ and $\pi(\vec{C}_{\varepsilon}^{U'}, p^U) = \pi(\vec{C}^{U'}, p^U)$, which implies: (i) H -types remain indifferent to their new contract and the new U -type contract; (ii) U -types get additional insurance; and (iii) the *additional* insurance is priced at their actuarially fair rate.
3. $V(\vec{C}_{\varepsilon}^{L'}, p^L) - V(\vec{C}_{\varepsilon}^{U'}, p^L) = V(\vec{C}^{L'}, p^L) - V(\vec{C}^{U'}, p^L)$ and $C_{\varepsilon,1}^{L'} = C_{\varepsilon,2}^{L'}$, so that (i) L -types are now offered full insurance, and (ii) the benefit to a U -type of becoming informed and learning that she is an L -type is unchanged.

The allocation $(\vec{c}_\varepsilon^{L'}, \vec{c}_\varepsilon^{U'}, c_\varepsilon^{H'})$ is incentive and information compatible whenever L is observable or verifiable. Moreover, by steps 2 and 3, the value of information to U -types is unchanged; it therefore remains consistent with equilibrium for the same fraction κ of individuals to be informed.

Note that $\pi(\vec{C}_{\varepsilon=0}^{L'}, p^L) > \pi(\vec{C}^{L'}, p^L)$, as $\vec{C}_{\varepsilon=0}^{L'}$ and $\vec{C}^{L'}$ lie on the same indifference curve and the former provides more insurance. Hence, the menu $(\vec{C}_{\varepsilon=0}^{L'}, \vec{C}_{\varepsilon=0}^{U'}, \vec{C}_{\varepsilon=0}^{H'})$ earns non-negative profits. By continuity, $(\vec{C}_{\varepsilon}^{L'}, \vec{C}_{\varepsilon}^{U'}, \vec{C}_{\varepsilon}^{H'})$ also earns non-negative profits for sufficiently small ε , and the menu $(\vec{C}_{\varepsilon}^{L'}, \vec{C}_{\varepsilon}^{U'}, \vec{C}_{\varepsilon}^{H'})$ is

therefore resource-feasible. Hence, it could in principle be implemented by a social planner no more informed than the market. Since it Pareto dominates the banned risk-classification market outcome $(\vec{C}^{L'}, \vec{C}^{U'}, \vec{C}^{H'})$, risk classification is inefficient whenever L is observable or verifiable—i.e., in the informational environments of sections 6.1 and 6.3.

The inefficiency of banning risk classification when information status but not risk type is observable, as in section 6.2, can be established using a related argument that relies on giving U - instead of L -types full insurance when the ban is lifted.²⁷

Suppose next that the market outcome when risk classification is banned involves $\kappa = 1$. The market outcome will therefore be the MWS outcome $(\vec{C}_{MWS}^L, \vec{C}_{MWS}^H)$ for a two-type economy. This is inefficient in the informational environments of section 6.1 and 6.3, since a social planner can feasibly implement the contract pair $(\vec{C}^{L*}, \vec{C}_{MWS}^H)$, where $\vec{C}^{L*}$ is the full insurance contract with $\pi(\vec{C}^{L*}, p^L) = \pi(\vec{C}_{MWS}^L, p^L)$, (as in Figure 1), which will remain consistent with $\kappa = 1$. It is also inefficient in the informational environment of section 6.2, since in this case, a social planner can instead offer the menu $(\vec{C}_{MWS}^L, \vec{C}^{U*}, \vec{C}_{MWS}^H)$, where $\vec{C}^{U*}$ is the full insurance contract satisfying $\pi(\vec{C}^{U*}, p^U) = \lambda\pi(\vec{C}_{MWS}^H, p^H) + (1 - \lambda)\pi(\vec{C}^{L*}, p^L)$. This will yield an equilibrium with $\kappa = \kappa_0$ since uninformed types will strictly prefer $\vec{C}^{U*}$ to paying τ to become informed and then face a lottery with the same expected value. This equilibrium thus Pareto dominates the original banned-classification equilibrium.

We conclude that the outcomes without risk classification are inefficient in each of the three information environments in which risk classification could potentially be used.

6.6 Distributional consequences of risk classification

As in Section 5, risk classification will often have distributional as well as efficiency consequences in these endogenous informational environments. For example, H -types may receive positive cross subsidies in a banned-classification pure private information market outcome like that depicted in Figure

²⁷ When testing status is unobservable, and it is (a) possible for individuals who have taken a test to verify having taken it, (b) impossible for people who have *not* been tested to verify it, and (c) firms cannot *ever* observe the *outcome* of any test, then one can construct examples where a ban on the use of test status has no negative efficiency consequences and beneficial distributional consequences. Specifically, If $\kappa_0\lambda$ and $(1 - \kappa_0)/(\kappa_0(1 - \lambda))$ are both sufficiently high, then the banned-classification equilibrium will involve H -types getting full, fair insurance and will involve U and L types *pooling* at a less-than full insurance allocation. Lifting the ban does not affect H -types but breaks the L -and- U pooling, undesirably making the former better off at the latter's expense. Since both types' feasible allocations remain constrained by the same H -type incentive compatibility constraint, however, "undoing" this redistribution and strictly improving on the original pooling allocation is impossible.

7. If so, banning risk classification will increase the well-being of H -types relative to the symmetric information market outcomes $\vec{C}^{i\#}$.

The welfare consequences of risk classification with endogenous private information are thus qualitatively quite similar to the consequences with fixed private information (per section 5). A ban on risk classification may lead to a market outcome which is preferable—on distributional grounds—to the market outcome which obtains when risk classification is permitted. But *in principle* there is a more efficient way to achieve these redistributive goals. Of course, as we discuss further below, whether it is possible *in practice* to achieve these goals in a more efficient way is far from clear: there is no obvious analog, e.g., to the partial social insurance scheme discussed in section 5.2 which allows the social planner to “lock in” the distributional benefits of a risk classification ban while removing it.

6.7 Robustness to other types of testing

The preceding analysis focused on a situation in which (a) there are no *intrinsic* benefits of testing and (b) testing is perfect. The basic conclusion that banning classification has efficiency consequences is robust to relaxing both assumptions.

The private benefits of testing in the preceding analysis, if any, arise from the informational asymmetries and the associated sorting of individuals into insurance contracts. In practice, of course, testing may also have direct benefits, for example by facilitating better treatments. A simple way of incorporating this sort of benefit is to allow the cost of testing τ to be negative. While this can change the market outcomes discussed in sections 6.1 and 6.2 (since individuals could now choose to get tested), it does not substantively affect the economic analysis of *risk classification* in these environments. Hoel et al. (2006) reach a similar conclusion in a model with heterogeneity in the benefits of testing.²⁸

Browne and Kamiya (2012) study the demand for underwriting. Specifically, they analyze a model of costly and potentially imperfect underwriting in a two-type market where all individuals are perfectly and, initially, privately informed about their type. By going through an underwriting process, L -types can signal their type to insurers and thereby improve their coverage. (H -types may also demand underwriting if the test is sufficiently imperfect that they find it worth bearing the underwriting cost for

²⁸ Also see Doherty and Posey (1998) and Hoel and Iversen (2000) who build on Hoy’s (1989) model of self-protection to study the consequences of genetic testing in a market with asymmetric information.

the chance at being incorrectly classified as a low risk.) Banning underwriting is generally inefficient in this context as well, by the basic argument in Rothschild (2011).

7. Endogenous Risk Classification

Endogenous risk classification refers to risk classification based on the *choices* made by the insured individuals rather than their intrinsic characteristics. These choices may affect individual's riskiness (and hence may be related to moral hazard) as in Bond and Crocker (1991) or they may simply *signal* riskiness, as in Polborn (2008).

In Bond and Crocker (1991), preferences are given by

$$V^i(C_1, C_2, x) = (1 - p^i(x))u(C_1) + p^i(x)u(C_2) + \theta^i G(x), \quad (11)$$

where x is an endogenous choice that has a direct effect on individual well-being via the strictly concave and increasing function $G(x)$ and an indirect effect through the accident probabilities $p^i(x)$. The function $G(x)$ could, for example, represent the “enjoyment” of smoking cigarettes. The parameter θ^i captures the i -type's intrinsic taste for x , which is assumed to have unit cost c . The accident risks $p^i(x)$ for the two types $i = H, L$ potentially differ for two reasons. First, it may be that $p^H(x) > p^L(x)$ for any given x , so that H -types are *intrinsically* riskier. Second, the H -type may have a stronger preference for the risky activity $\theta^H > \theta^L$ and may therefore choose a larger x . For expositional simplicity, we focus here on the second effect by taking $p^H(x) = p^L(x) \equiv p(x)$ so that individuals differ only in their taste for x not in their intrinsic riskiness. If x represents cigarette smoking, this amounts to assuming that the health risks associated with a given level of cigarette smoking are independent of how much an individual intrinsically enjoys smoking.

“Endogenous risk classification” in this context involves pricing policies based on the observable decision x . It is straightforward to show that endogenous risk classification implies a first-best efficient market outcome (when $p^H(x) = p^L(x)$). To wit: a first best allocation maximizes $V^i(C_1^i, C_2^i, x^i)$ subject to a resource constraint

$$[p(x)((W - cx) - D) + (1 - p(x))(W - cx)] - [p(x)C_2^i + (1 - p(x))C_1^i] \geq 0. \quad (12)$$

which can equivalently be interpreted as a “break-even” or zero-profit constraint for insurers. The first order necessary conditions require first that $C_1^i = C_2^i \equiv C^{i*}$, i.e., full insurance given x , and, second, that x is chosen so that the iso-resource set (a constant left-hand-side of (12)) and the i -type’s indifference set are tangent in (x, C_1, C_2) -space.

Figure 9 plots the tangency conditions for $i = H, L$, making use of the full insurance property to reduce the dimensionality and plot in (x, C^*) space. Because we assume $p^H(x) = p^L(x)$, the zero-profit line is the same for both types. (It is non-linear because $p(x)$ depends on x .) The first-best allocation occurs at the tangency points of the indifference curves (labeled IC^H and IC^L). Because $\theta^H > \theta^L$, the tangency points (x^{*H}, C^{*H}) and (x^{*L}, C^{*L}) satisfy $x^{*H} > x^{*L}$. It is clear from Figure 9 that this allocation is incentive compatible. It is therefore the unique Nash equilibrium of the Rothschild-Stiglitz style game where profit-maximizing firms simultaneously offer contracts and individuals then choose contracts. This means that endogenous classification—here captured by the fact that firms are offering different contracts to individuals who make different choices of x —leads to a first-best efficient outcome.

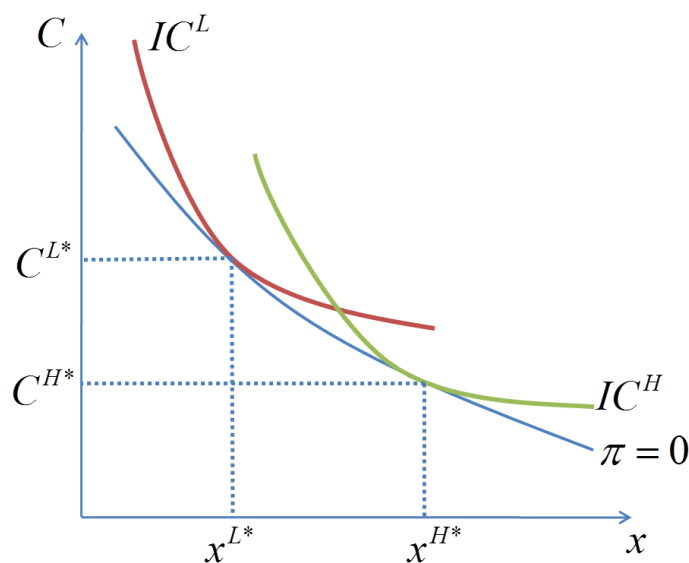


Figure 9: First-Best Break-Even Allocations with Endogenous Risk Classification

If classification is banned, so that contracts are not allowed to condition on x , on the other hand, no first-best allocation is feasible. Intuitively, a ban on x -based classification introduces an inefficiency

causing moral hazard problem: individuals do not take the effects of x on $p(x)$ into account when they are fully insured and their contract is independent of x .²⁹

The preceding argument for the inefficiency of banning endogenous risk classification in Bond and Crocker's (1991) environment is clearly robust to small differences between $p^H(\cdot)$ and $p^L(\cdot)$, which would appear in Figure 9 as vertical differences between the break-even constraints of the two types. The argument depends critically, however, on firms having significant flexibility in designing their insurance products—in particular, in their ability to offer contracts whose payouts depend on the choice x (or simply the ability to dictate x).

Polborn (2008) studies a loosely related model of endogenous risk classification without this flexibility: his model has a fixed contract and mandated purchases. Endogenous risk classification arises because risk type is correlated with preferences for something observable, such as car color. In this context, the *use* of risk classification—pricing based on endogenously chosen observables—can cause inefficiencies. Intuitively, this is because observable-based competitive pricing leads to lower prices for the observables chosen by individuals with low risk. Individuals' choices in equilibrium are thus inefficiently distorted towards those observables that are intrinsically preferred by these low-risk individuals.

The contrast between the efficiency properties of risk classification in Bond and Crocker's and Polborn's models of endogenous risk classification thus mirror the contrast between the efficiency properties of risk classification in Section 3 and Section 5: when firms and consumers have more “flexibility” to respond to restrictions on risk classification, there will typically be more inefficiency. Intuitively: inefficiencies are caused by the attempt to avoid or “work around” restrictions on risk classification.

8. New Directions

Our analysis of the economic consequences of risk classification has focused on the class of relatively simple models discussed most extensively in the literature, namely those with: (1) a small number of types (and signals); (2) adverse selection but no (or only a limited form of) moral hazard; (3) static rather

²⁹ Formally: the first order conditions for a first-best x is $\theta G'(x^*) - cu'(C^*) = Dp'(x^*)$. The utility an individual gets from choosing x and a full insurance contract at a premium R is $\tilde{U}(x) \equiv u(W - R - cx) + \theta G(x)$, the derivative of which is $\theta G'(x) - cu'(W - R - cx)$. Since $Dp'(x^*) > 0$, individuals find a marginal increase in x desirable at any first best optimum allocation.

than dynamic contracting; (4) exclusive contracting and non-linear pricing; and (5) a single risk. In this section, we briefly discuss the less developed literatures which relax these assumptions.

8.1 Models with richer type spaces

Two-type settings have the advantage of facilitating tractable, graphical analyses. The qualitative insights from two-type models often generalize; but not always. Within the two-type symmetric information framework discussed above, for example, bans on risk classification can have negative efficiency consequences in both the “fixed-full contracts” case (where firms compete on the price of full-insurance contracts) and the “screening” case (where firms compete on both price and coverage levels). But the inefficiencies in these two cases are qualitatively distinct: With fixed contracts, inefficiencies arise when pooled pricing causes market unraveling and low-risk types exit the market entirely. With screening, the low-risk type remains insured but is inefficiently quantity-rationed. Recent work by Hendren (2013) shows that this qualitative distinction is not general. In particular, his “No Trade” theorem shows that *unraveling* can occur even when firms can employ screening mechanisms.

Specifically, Hendren extends the canonical one-risk insurance market framework described above to allow for an *arbitrary* distribution $F(p)$ of types differing in their risk $p \in [0,1]$ of experiencing a loss. He shows that an insurance market will completely unravel—there will be No Trade—*unless* there is some risk type p^* who is willing to pay the average cost of insurance for the pool of all higher risk types in order to obtain a small quantity of insurance. Figure 10 illustrates a situation in which the No-Trade condition is violated. To wit: the slope of the p^* -type’s indifference curve at the endowment point is lower than the slope of the zero-profit curve for the pooled above- p^* risk types, namely $(1 - \mathbb{E}_\Lambda[p|p \geq p^*])/\mathbb{E}_\Lambda[p|p \geq p^*]$, where $\mathbb{E}_\Lambda$ is the expectation-with-respect-to-measure- Λ operator. . So the p^* type *does* value insurance enough to buy it even when it is priced fairly for the pooled set of risk types above p^* . When there is no such type, however,³⁰ the only resource feasible, incentive compatible, and individually rational allocation is the autarkic allocation with no insurance at all.

³⁰ For example, when $u(C) = \ln(C)$, $\Lambda(p)$ is uniform on $[0,1]$, and $D \leq \frac{W}{2}$.

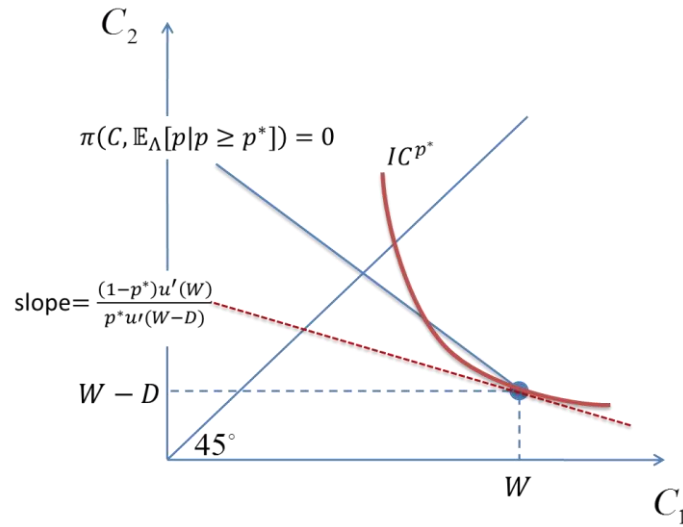


Figure 10: Hendren's "No-Trade" Condition is Violated at p^*

Hendren's No-Trade theorem has important implications for understanding the consequences of risk classification. First, it directly implies that eliminating risk classification in a symmetric information world in which the No-Trade condition is satisfied will lead to complete market unraveling—and this is true regardless of whether firms can, in principle, employ screening mechanisms.³¹

The No-Trade theorem also has important implications for understanding the consequences of banning *partially* informative risk classification. Consider a market with a full type support $p \in [0,1]$ and suppose that there is a signal $\sigma \in \{A, B\}$ that is correlated with risk. Assume, in particular, that B is the riskier category, in the sense that $\mathbb{E}_{\Lambda_B}[p|p \geq p^*] \geq \mathbb{E}_{\Lambda_A}[p|p \geq p^*]$ for all p^* . Then there are four possibilities. (1) If the No-Trade condition holds *within* the A category, then it will, *a fortiori*, also hold for the population as a whole and within the B -category. A ban on risk classification is then irrelevant: there will be no trade both with and without a ban. (2) If the No-Trade condition is violated within the A category but is satisfied for the pooled population, then there will be no market for insurance under a ban on risk classification, while removing the permitting risk classification can³² lead to welfare-improving insurance provision—although only *within* category A . (3a) If the No-Trade condition is violated within the pooled population but not within category- B , then trade can take place under a ban on risk classification and

³¹ In fact, as Hendren (2014) shows, the No-Trade theorem implies that the very structure of the two-type Rothschild-Stiglitz (1976) *equilibrium* is not robust: if the type distribution Λ has full support near $p = 1$, then the Rothschild-Stiglitz equilibrium exists only if it involves no insurance for any type. As such, its prediction that some types will be fully insured and other types will be partially insured holds only when the risk is essentially bounded away from $p = 1$. This is among the reasons we focus on the MWS equilibrium concept here.

³² There is *scope* for trade in the A category, and there is always a non-trivial MWS equilibrium with trade when the No Trade condition is violated (*viz* Hendren, 2014, footnote 11).

removing the ban can eliminate category- B individuals' access to insurance. (3b) If the No-Trade condition is violated within category B , then (different) trade can take place with and without a ban.

This typology is qualitatively similar to the typology of section 5.1 with “the No-Trade condition is violated” here replacing “scope for Pareto improving cross subsidies” in section 5.1. A similar efficiency-distribution tradeoff typology applies as well: in case (1), bans on risk-classification are irrelevant. In case (2), bans on risk-classification have purely negative efficiency consequences without any beneficial redistribution. In cases (3a) and (3b) there are (at least potentially) distributional consequences of a ban, with the worse-off category B individuals benefiting from a ban at the expense of the better off category A individuals. The big difference from case (3) in section 5.1 is that in case (3a) here, permitting risk classification doesn't just make the B -category types worse off, but in fact completely eliminates their access to insurance.³³

Perhaps less obviously, there are also generically negative efficiency consequences of a risk-classification ban in cases (3a) and (3b)—at least under the MWS market equilibrium concept. To see this, suppose that in the presence of a ban the market implements some constrained efficient allocation that involves some implicit per-capita cross-subsidies from the-risk category A to the high-risk category B individuals. Let R_σ denote the expected value of profits in category σ (so $R_A > 0 > R_B$) and let $\bar{V}(p)$ denote the expected utility of the p -type in this allocation. Then consider the (mathematical) problem of maximizing for each $\sigma \in \{A, B\}$ the average expected utility (e.g.) of the category- σ individuals subject to incentive compatibility, a resource constraint that the per-capita profits are no more than R_σ , and a set of minimum utility constraints $V(p) \geq \bar{V}(p)$. These problems are feasible since the original banned-classification equilibrium satisfies these constraints. The minimum utility constraints ensure that no individual is worse off in their solutions than in the banned categorization equilibrium. Since there is increased scope for Pareto improving cross-subsidies within the lower-risk category A , however, the solution to the category- A problem will in general Pareto dominate the banned-classification equilibrium. In other words, as in case (3) in section 5.1, the banned-classification equilibrium will be constrained inefficient.

In cases (3a) and (3b), bans on risk classification thus involve tradeoffs between efficiency and distributional goals. As in the endogenous information environments discussed in section 6.5, the case for focusing exclusively on the inefficiencies is significantly less compelling in this setting than in section

³³ Hendren (2014) provides compelling evidence that firms *do* use risk classification to *exclude* individuals from insurance purchases, and do so precisely among classes where the No-Trade condition is most likely to be satisfied.

5.2: there is *theoretical* scope for Pareto improving on a banned-classification equilibrium in this case, but it is not obvious that there is also a *practical* way of implementing these improvements.

8.2 Moral hazard extensions

Bond and Crocker's (1991) endogenous risk classification model becomes a "private action" moral hazard setting when endogenous risk classification is banned, but it is not, at its core, about risk-classification with moral hazard, since risk-classification in their setting is equivalent to *eliminating* the moral hazard problem. An interesting and important direction for future work would be to explore the consequences of restricting classification based on observable *correlates* of propensities towards unobservable but risk-relevant actions. This is a challenging problem, as screening models with both adverse selection and moral hazard are notoriously difficult to solve—since, as discussed in Arnott (1992), Chassagnon and Chiappori (1997), and Stiglitz and Yun (2013), the introduction of moral hazard typically leads to single-crossing violations.

One potentially promising route around these technical challenges is to abstract from "screening" considerations. Einav et al. (2013), for example, study and empirically estimate a model with two types of selection: traditional adverse selection based on risk, *and* "selection on moral hazard," whereby privately known heterogeneity in the responsiveness of treatment intensity to coverage leads to different preferences over an exogenously fixed set of contracts.

8.3 Dynamic extensions

The preceding analysis is relevant to—but unsuitable for a complete analysis of—risk classification based on new information that arises over time as individuals (and/or firms) learn more about their riskiness. If, for example, health insurance contracts are negotiated annually and firms can use pre-existing conditions to classify risks, then developing an expensive-to-treat, chronic condition in one year will raise an individual's premiums or limit her access to insurance in future years. From an *ex-ante* perspective, then, permitting risk classification based on pre-existing conditions *in future years* exposes individuals to significant welfare-reducing reclassification risk. A static model is obviously insufficient for fully analyzing this dynamic phenomenon; nevertheless, in any given future period, the conclusion from static models that restricting risk classification can have efficiency costs still applies. Risk-classification in dynamic settings can thus involve tradeoffs between *ex-ante* efficient provision of insurance against reclassification risk and *interim* efficient insurance provision.

Various types of dynamic contracts have been suggested in the literature as ways to ease these tradeoffs. Long-term contracts with full two-sided commitment can typically provide full insurance against both types of risk, but the literature has focused on the more realistic case with limited commitment on the insurance buyer's side. Tabarrok (1994) suggests developing markets for insuring the re-classification risk (specifically, the risk arising from genetic tests) directly. Pauly et al. (1995) suggest guaranteed renewable contracts and show that these can be used to *fully* insure reclassification risk in symmetric information environments.³⁴ Hendel and Lizzeri (2003) show,³⁵ however, that simple guaranteed renewable contracts are *not* optimal for individuals who face capital market imperfections: optimal contracting involves tradeoffs between classification risk insurance and consumption smoothing motives. Since classification risk is not fully eliminated by the private market, interventions that limit risk-classification are potentially efficient.

Polborn et al. (2006) identify a similar tradeoff in a dynamic model of life insurance. In their model, individuals learn over time about *both* their risk type and their life insurance needs. Uncertainty about life insurance needs militates for delaying the purchase of life insurance, but these delays expose individuals to classification risk—and, again, restrictions on risk classification *may* improve welfare.

8.4 Non-exclusive contracting, linear pricing, and multi-state models

We have so far focused on exclusive contracting models. Some markets, notably those for life insurance and life annuities, are non-exclusive: individuals can, and do, purchase simultaneous policies from multiple insurers. The literature typically assumes linear pricing when modeling these environments.³⁶ Hoy (2006) discusses the effect of banning risk classification in a two-type asymmetric information model with a single risk and with non-exclusivity-*cum*-linear pricing. His analysis indicates that the welfare impacts of such bans are qualitatively similar to those in the analogous exclusive contracting framework: there are negative efficiency consequences of these bans, but they can also have distributional effects that lead to an increase in (e.g.) utilitarian social welfare.

We have also focused on models with a single insurable risk; this is consistent with the literature and reflects the fact that single-risk models are analytically simpler and typically yield qualitative predictions

³⁴ Pauly et al. (2011) show that guaranteed renewability can still provide full insurance against reclassification risk even with substantial informational asymmetries

³⁵ This builds on a long line of literature on dynamic contracting with limited commitment, including Dionne and Doherty (1994).

³⁶ *Viz* Abel (1986), Villeneuve (2003), Hoy and Polborn (2000), and Polborn et al. (2006).

similar to many-risk models. In non-exclusive contracting environments with linear pricing, however, distinguishing between single-risk and many-risks is qualitatively important. As Brunner and Pech (2005) and Rothschild (2014) point out—both in the annuity market context—linear pricing in a single risk model implies *pooling*, since all types purchase the same contract. When there are multiple payout states (future payment periods in the annuity context), separation via different across-state payment patterns is possible—and, indeed, Finkelstein and Poterba (2002, 2004) provide evidence of exactly this sort of risk-based separation. The implications of restrictions on risk-based classification with linear pricing and contract-shape-based screening have not yet been studied comprehensively.

9. Conclusions

We have analyzed the consequences of restrictions on risk classification in a broad range of canonical insurance market models. Such restrictions have potentially desirable distributional consequences; indeed, that is a major motivation for imposing such restrictions. We argued that such restrictions typically also have negative efficiency consequences in market-based settings that are otherwise unregulated. These negative efficiency consequences mean that, in principle, there is some method for achieving the distributional benefits of such restrictions at a lower cost without imposing such restrictions. Insofar as it is possible *in practice* to obtain these distributional benefits in lower-cost ways—as in the settings discussed in section 5.2—these negative efficiency consequences argue strongly against restricting the use of risk-classification in otherwise market-based settings. There are, of course, alternative interpretations of this result: one can interpret it as an explicit “pro-market” argument *against* bans on risk classification, or, alternatively as a “pro-interventionist” argument for the implementation of alternatives or complements to bans on risk classification.

In other, and perhaps most, settings (viz Sections 6.6 and 8.1), there is no clear *practical* alternative to risk classification bans, for obtaining its (potential) distributional benefits at lower efficiency costs. In such settings, alternative methods for evaluating its efficiency and equity consequences are necessary. As such, the practical feasibility, and, indeed, the practical design of these *alternative* ways of achieving the distributional benefits of categorical pricing restrictions should be a significant component of any practical evaluation of risk-classification policies. Our hope is that laying out the economic consequences

of risk classification in a reasonably comprehensive set of environments, as we have done in this article, will facilitate such evaluations in the future.³⁷

References

- Abel, A. B. (1986). Capital Accumulation and Uncertain Lifetimes with Adverse Selection. *Econometrica* **54**, 1079–1097.
- Akerlof, G.A. (1970). The market for “lemons”: Quality uncertainty and the market mechanism. *The Quarterly Journal of Economics* **84**, 488-500.
- American Academy of Actuaries (2001). Risk classification in voluntary individual disability income and long-term care insurance. http://www.actuary.org/pdf/health/issue_genetic_021601.pdf.
- Arnott, R.J. (1992). Moral hazard and competitive insurance markets. In G. Dionne, ed., *Contributions to Insurance Economics*. 325-358. Norwell, MA: Kluwer Academic Publishers.
- Arrow, K.J. (1963). Uncertainty and the welfare economics of medical care. *The American Economic Review* **53**, 941–973.
- Baker, T. (2011). Health insurance, risk, and responsibility after the patient protection and affordable care act. Research paper 11-03. ILE, University of Pennsylvania Law School.
- Bond, E. and Crocker, K. (1991). Smoking, skydiving, and knitting: The endogenous categorization of risks in insurance markets with asymmetric information. *Journal of Political Economy* **99**, 177-200.
- Browne, M. J., & Kamiya, S. (2012). A Theory of the Demand for Underwriting. *Journal of Risk and Insurance* **79**, 335-349.
- Brunner, J.K., and Pech, S. (2005). Adverse Selection in the Annuity Market when Payoffs Vary over the Time of Retirement. *Journal of Institutional and Theoretical Economics* **161**, 155–183.
- Buchmueller, T. and DiNardo, J. (2002). Did community rating induce an adverse selection death spiral? Evidence from New York, Pennsylvania, and Connecticut. *American Economic Review* **92**, 280-294.

³⁷ While we have largely ignored them in this essay, we also recognize that “non-economic” concerns, for example about the *intrinsic* unfairness of discriminatory risk classification are also an important component of real-world policy considerations (viz, e.g., Thiery and Van Schoubroeck (2006) and Thomas (2007)) and should not be ignored in evaluating risk classification policies.

- Buzzacchi, L. and Valletti, T. M. (2005). Strategic price discrimination in compulsory insurance markets. *The Geneva Risk and Insurance Review* **30**, 71-97.
- Chade, H., and Schlee, E. (2012). Optimal insurance with adverse selection. *Theoretical Economics* **7**, 571-607.
- Chassagnon, A., and Chiappori, P. A. (1997). Insurance under moral hazard and adverse selection: The case of pure competition. *DELTA-CREST Document*.
- Chiappori, P.A. (2006). The welfare effects of predictive medicine. In Chiappori, P.A. & Gollier, C. (eds.) *Competitive failures in insurance markets: Theory and policy implications*. 55-79. Cambridge, MA: MIT Press.
- Crocker, K.J. and Snow, A. (1985a). The efficiency of competitive equilibria in insurance markets with asymmetric information. *Journal of Public Economics* **26**, 207-219.
- Crocker, K.J. and Snow, A. (1985b). A simple tax structure for competitive equilibrium and redistribution in insurance markets with asymmetric information. *Southern Economic Journal* **51**, 1142-1150.
- Crocker, K.J. and Snow, A. (1986). The Efficiency Effects of Categorical Discrimination in the Insurance Industry. *Journal of Political Economy* **94**, 321-344.
- Crocker K. and Snow, A. (1992). The social value of hidden information in adverse selection economies. *Journal of Public Economics* **48**, 317-347.
- Crocker, K. and Snow, A. (2013). The theory of risk classification. In Dionne (ed.), *Handbook of Insurance*. 281-314. New York: Springer.
- Cutler, D.M. and Reber, S.J. (1998). Paying for health insurance: the trade-off between competition and adverse selection. *The Quarterly Journal of Economics*, **113**, 433-466.
- Dionne, G. and Doherty, N.A. (1994). Adverse selection, commitment and renegotiation: Extension to and evidence from insurance markets. *Journal of Political Economy* **102**, 210-235.
- Dionne, G., Fombaron, N. and Doherty, N.A. (2013). Adverse selection in insurance contracting. In Dionne, G. (ed.), *Handbook of Insurance* 231-280. New York: Springer.
- Doherty, N. A. and Posey, L.L. (1998). On the value of a checkup: Adverse selection, moral hazard and the value of information. *Journal of Risk and Insurance* **65**, 189-211.
- Doherty, N.A. and Thistle, P.D. (1996). Adverse selection with endogenous information in insurance markets. *Journal of Public Economics* **63**, 83-102.
- Dubey, P. and Geanakoplos, J. (2002). Competitive pooling: Rothschild-Stiglitz reconsidered. *The Quarterly Journal of Economics* **117**, 1529-1570.

- Einav, L. and Finkelstein, A. (2011). Selection in Insurance markets: Theory and empirics in pictures. *Journal of Economic Perspectives* **26**, 1-17.
- Einav, L., Finkelstein, A., Ryan, S., Schrimpf, P., and Cullen, M. (2013). Selection on Moral Hazard in Insurance Markets. *American Economic Review* **103**, 178-219.
- Finkelstein, A. and Poterba, J. (2002). Selection Effects in the Market for Individual Annuities: New Evidence from the United Kingdom. *Economic Journal* **112**, 28–50.
- Finkelstein, A., and Poterba, J. (2004). Adverse Selection in Insurance Markets: Policyholder Evidence from the U.K. Annuity Market. *Journal of Political Economy* **112**, 183–208.
- Finkelstein, A., Poterba, J. and Rothschild, C. (2009). Redistribution by insurance market regulation: Analyzing a ban on gender-based retirement annuities. *Journal of Financial Economics* **91**, 38-58.
- Harrington, S. (2010a). The health insurance reform debate. *Journal of Risk and Insurance* **77**, 1, 5-38.
- Harrington, S. (2010b). U.S. health-care reform: The patient protection and affordable care act. *Journal of Risk and Insurance* **77**, 3, 703-708.
- Harsanyi, J.C. (1953). Cardinal utility in welfare economics and in the theory of risk-taking. *Journal of Political Economy* **61**, 434-435.
- Harsanyi, J.C. (1955). Cardinal welfare, individualistic ethics, and interpersonal comparisons of utility. *Journal of Political Economy* **63**, 309-321.
- Hellwig, M.F., 1987. Some recent developments in the theory of competition in markets with adverse selection. *European Economic Review* **31**, 319-325.
- Hendel, I. and Lizzeri, A. (2003). The role of commitment in dynamic contracts: Evidence from life insurance. *Quarterly Journal of Economics* **118**, 299–327.
- Hendren, N. (2013). Private information and insurance rejections. *Econometrica* **81**, 1713-1762.
- Hendren, N. (2014). Unraveling versus Unraveling: A Memo on Competitive Equilibriums and Trade in Insurance Markets. THIS ISSUE
- Hoel, M., Iversen, T., Nilssen, T. and Vislie, J. (2006) Genetic testing in competitive insurance markets with repulsion from chance: A welfare analysis. *Journal of Health Economics* **26**, 251-262.
- Hoel, M. and Iversen, T. (2000). Genetic testing when there is a mix of compulsory and voluntary health insurance. *Journal of Health Economics* **21**, 253-270.
- Hoy, M. (1982). Categorizing risks in the insurance industry. *Quarterly Journal of Economics* **97**, 321-336.
- Hoy, M. (1989). The value of screening mechanisms under alternative insurance possibilities. *Journal of Public Economics* **39**, 177-206.

- Hoy, M. (2006). Risk classification and social welfare. *The Geneva Papers* **31**, 245–269.
- Hoy, M. and Polborn, M. (2000). The value of genetic information in the life insurance market. *Journal of Public Economics* **78**, 235-252.
- Hoy, M. and Witt, J. (2007). Welfare effects of banning genetic information in the life insurance market: The case of BRCA1/2 Genes. *Journal of Risk and Insurance* **74**, 523–546.
- Martin, A. (2007). On Rothschild-Stiglitz as competitive pooling. *Economic Theory* **31**, 371-386.
- Mimra, W. and Wambach, A. (2011). A game-theoretic foundation for the Wilson equilibrium in competitive insurance markets with adverse selection, CESifo working paper: Industrial Organisation, No. 3412.
- Miyazaki, H. (1977). The rat race and internal labor markets. *Bell Journal of Economics* **8**, 394-418.
- Morrissey, M. (2013). Health Insurance in the United States. In Dionne, G. (ed.), *Handbook of Insurance*. 957–996. New York: Springer.
- Netzer, N. and Scheuer, F. (2013). A Game Theoretic Foundation of Competitive Equilibria with Adverse Selection. *International Economic Review*. Forthcoming.
- Pauly, M.V. (1974). Overinsurance and public provision of insurance: The roles of moral hazard and adverse selection. *Quarterly Journal of Economics* **88**, 44-62.
- Pauly, M. V., Kunreuther, H., and Hirth, R. (1995). Guaranteed renewability in insurance. *Journal of Risk and Uncertainty* **10**, 143-156.
- Pauly, M. V., Menzel, K., Kunreuther, H., and Hirth, R. A. (2011). Guaranteed renewability uniquely prevents adverse selection in individual health insurance. *Journal of Risk and Uncertainty* **43**, 127-139.
- Picard, P. (2014). Participating insurance contracts and the Rothschild-Stiglitz equilibrium puzzle. THIS ISSUE.
- Polborn, M., Hoy, M., and Sadanand, A. (2006) Advantageous effects of adverse selection in the life insurance market. *The Economic Journal* **116**, 327-354.
- Polborn, M. (2008). Endogenous categorization in insurance. *Journal of Public Economic Theory* **10**, 1095–1113.
- Riley, J.G., (1979). Informational Equilibrium. *Econometrica* **47**, 331-359.
- Rothschild, C. (2011). The efficiency of categorical discrimination in insurance markets. *Journal of Risk and Insurance*, **78**, 267–285.
- Rothschild, C. (2014). Nonexclusivity, linear pricing, and annuity market screening. *Journal of Risk and Insurance* doi: 10.1111/jori.12019

Rothschild, M and Stiglitz, J. (1976). Equilibrium in competitive insurance markets: An essay on the economics of imperfect information. *The Quarterly Journal of Economics* **90**, 629-649.

Spence, M. (1978). Product differentiation and performance in insurance markets. *Journal of Public Economics* **10**, 427-447.

Stiglitz, J. (1977). Monopoly, non-linear pricing and imperfect information: The insurance market. *Review of Economic Studies* **44**, 407-430.

Stiglitz, J., and Yun, J. (2013). Optimality and Equilibrium In a Competitive Insurance Market Under Adverse Selection and Moral Hazard. National Bureau of Economic Research working paper 19317.

Strohmenger R. and Wambach, A. (2000). Adverse selection and categorical discrimination in the health insurance market: the effects of genetic tests. *Journal of Health Economics* **19**, 197-218.

Tabarrok, A. (1994). Genetic testing: an economic and contractarian analysis. *Journal of Health Economics* **13**, 75-91.

Thiery, Y., and Van Schoubroeck, C. (2006). Fairness and Equality in Insurance Classification. *The Geneva Papers on Risk and Insurance-Issues and Practice* **31**, 190-211.

Thomas, R. G. (2007). Some novel perspectives on risk classification. *The Geneva Papers on Risk and Insurance-Issues and Practice* **32**, 105-132.

Thomas, R. G. (2008). Loss coverage as a public policy objective for risk classification schemes. *Journal of Risk and Insurance*, 75(4), 997-1018.

Villeneuve, B. (2003). Mandatory pensions and the intensity of adverse selection in life insurance markets. *The Journal of Risk and Insurance* **70**, 527-548.

Villeneuve, B. (2005). Competition between insurers with superior information. *European Economic Review* **49**, 321-340.

Wilson, C. (1977). A model of insurance markets with incomplete information. *Journal of Economic Theory* **16**, 167-207.