

Models of Selection Markets

by

Christopher Wu

A thesis submitted in partial fulfillment of the requirements for the degree of

Master of Science

in

Applied Mathematics

Department of Mathematical and Statistical Sciences

University of Alberta

© Christopher Wu, 2024

Abstract

Selection markets describe markets in which people strategically “select” into certain options based on knowledge only they possess, and in doing so may communicate some of that knowledge. Examples include sick people buying more comprehensive health insurance contracts, talented students choosing more challenging programs, and high-quality goods producers offering longer warranties. These markets tend to be inefficient—sometimes to the point of total collapse. The theory of selection markets can help to diagnose issues and offer policy solutions. This thesis analyzes three theoretical problems in the theory of selection markets. The first part is a technical contribution connecting the recent Azevedo-Gottlieb model to the older reactive equilibrium model. The second part studies insurance markets with a fixed cost of providing contracts, and gives a necessary and sufficient condition under which the market will collapse. The third provides a framework for simulating selection markets with choice frictions, and constructs a model within this framework to study the ambiguous effect of rational inattention on insurance market inefficiency.

Acknowledgements

I would first and foremost like to thank my supervisor, Chris Frei, who has carried me through this process. My appreciations also go to Corinne Langinier and Brendan Pass for serving on my defence committee and providing helpful comments. An extended thanks goes to the mathematical finance and probability theory groups—both faculty and classmates.

Thanks also to my family for their support and to the teachers and professors who have mentored me over the years, but especially Jen Murphy, Kelvin Fong, Jik Chin, and Carlo Muro.

The research of Section 4 was enabled by the *Cedar* cluster, provided by the Digital Research Alliance of Canada (alliancecan.ca).

Table of Contents

1	Introduction	1
1.1	Background	1
1.2	Recent Advances and Challenges	3
1.3	Overview of This Thesis	5
2	Reconciling Azevedo and Gottlieb’s Equilibrium with Riley’s Initial Value Problem	7
2.1	Introduction	7
2.2	Model	10
2.3	On the Generality of the Model	11
2.4	Results	13
2.5	Discussion and Concluding Remarks	17
2.6	Chapter Appendix: Omitted Proofs	18
3	Fixed Costs, the Extensive Margin, and No-Trade	25
3.1	Introduction	25
3.2	The Riley Model	27
3.2.1	Assumptions	27
3.2.2	Equilibrium	29
3.3	D1 Equilibrium	31
3.4	Adding the Extensive Margin	32
3.5	Concluding Remarks	33
3.6	Chapter Appendix: Omitted Proofs	33

4	Dynamic Stochastic Selection Markets and the Implications of Inattention	39
4.1	Introduction	39
4.2	Background and Related Literature	41
4.3	Theoretical Framework	44
4.3.1	Basic Setup	44
4.3.2	Linear Logit Model	46
4.3.3	Discussion	47
4.4	Job Market Signaling	51
4.4.1	Setup	51
4.4.2	Theory	51
4.4.3	Results and Discussion	53
4.4.4	Relationship to the Intuitive “Speeches”	59
4.5	Implications for Insurance Markets	63
4.5.1	Introduction: A “Model-Free” Analysis	63
4.5.2	The Akerlof Model	65
4.5.3	The Riley-Rothschild-Stiglitz Model	69
4.6	A Health Insurance Model	70
4.6.1	Setup	71
4.6.2	Results	72
4.7	Concluding Remarks	75
4.8	Chapter Appendix	75
4.8.1	Models in the Limit	75
4.8.2	Health Insurance Model Simulation Details	78
4.8.3	Rational Inattention Foundations of the Utility Function	82
5	Conclusion	88

List of Figures

4.1	Riley outcome of Spence's signaling model	52
4.2	The "Noisy D1" in the equidistributed model	54
4.3	Convergence in the equidistributed model	54
4.4	Limiting distributions as a function of β after 10,000 periods	56
4.5	Marginal distributions for the case $\beta = 4$ after 25,000 periods	57
4.6	Signal distribution given $\mu(\theta_1)$ after 200,000 periods. All but 1/89 have converged.	58
4.7	Convergence of the case $\mu(\theta_1) = 1/89$ when the initial distribution is the D1 equilibrium	59
4.8	Akerlof unraveling of the market for high coverage contracts	66
4.9	A change in β causing an ambiguous effect	67
4.10	Example: Market for an infinitesimal amount of insurance	69
4.11	Distribution of coverage choices given various choices of $\log_{10}(\beta)$	73
4.12	$-\log_{10}(-W)$ given six choices of $\log_{10}(\beta)$	74

List of Symbols

$\mathcal{B}_Y$	Borel σ -algebra on a set Y
$(\Omega, \mathcal{F}, P)$	Probability space
$\frac{\partial f}{\partial x}$	Partial derivative of f with respect to x
$\int f(y) dy$	Lebesgue integral of f ; unless otherwise specified, the integral is with respect to the Lebesgue measure over the domain of f .
κ	Fixed cost
$\overline{\mathbb{R}}$	Extended reals $[-\infty, \infty]$
$\mathbb{R}_+$	Non-negative reals
$\subset$	Subset, not necessarily proper: I use $\subsetneq$ for strictly proper
θ, Θ	<i>Type</i> and <i>type space</i>
$\rightarrow, \uparrow, \downarrow$	Converges (pointwise unless otherwise specified), increases to, decreases to
$AC[a, b]$	The space of absolutely continuous functions on $[a, b]$
$c(x, \theta)$	Cost of contract x when purchased by type θ
$C^k(a, b)$	The space of k times continuously differentiable functions on (a, b)
E	Expected value operator

$E_{\eta \sim \mu(\xi, \cdot)}$	Expected value operator, with the respect to the probability measure $\mu(\xi, \cdot)$, where μ is a regular conditional probability of η given ξ
$f'(x)$ or $\frac{df}{dx}$	Derivative of f with respect to x . Note if α is not a function, α' usually denotes an alternative value of α
$f_x(x, y)$	Derivative of f with respect to the first variable
p	Price
$q^{y,p}(\cdot; \theta)$	Indifference curve of θ passing through (y, p)
$u(k)$	Utility function given consumption k
$U(x, p, \theta)$	Utility function of type θ buying contract x at price p
$x()$	<i>Allocation</i> , a map from type space Θ to contract space X ; sometimes called a <i>strategy</i>
x, X	<i>Contract</i> and <i>contract space</i> , in the abstract sense; sometimes called a <i>signal</i> , <i>message</i> , or simply an <i>action</i>

Chapter 1

Introduction

1.1 Background

Most markets involve a certain degree of information asymmetry between buyers and sellers. The seminal work of Akerlof (1970) showed that asymmetries can present a barrier to trade, if not completely collapse the market. If trade occurs at all, the market may be inefficient, with the informed side of the market making costly choices to convey information to the uninformed side. As such, the theory of markets with private information informs the debate on the role of government in these markets.

This thesis traces its roots to two canonical models. In the job market signaling model of Spence (1973), workers privately know their intrinsic productivity but cannot convey it directly to employers; in lieu, productive workers obtain education—a costly task which productive workers are willing to do in equilibrium but unproductive workers are not—to *signal* their productivity to potential employers. It is said here that productive workers *separate* from the *pool*. In the insurance market model of Rothschild and Stiglitz (1976), insurers *screen* potential buyers—who privately know their risk level—by offering contracts with different levels of coverage; those who are lower risk will be more willing to accept lower coverage, suboptimal as it may be, to separate themselves from higher risks and gain access to lower premiums.

A general conundrum which the authors produce is that there are cases with no Nash equilibrium, where any pool of insureds could have its best “cream skimmed” away (i.e. separated) by firms offering lower coverage at lower premiums, and any separating state—in which low-risk people buy low coverage and high-risk people buy high coverage—can be disrupted by offering a pool which is strictly preferable for both, in which high-risk people get better prices and low-risk people get more coverage.

We refer to the object being purchased here—labour, insurance, etc.—simply as a *contract*, which encodes all relevant non-price information—level of education, coverage, etc. The distinguishing feature of these markets is that price is determined not only by the contract, but by who *selects* into the contract. This thesis refers to such markets simply as *selection markets*.

The textbook treatment of signaling and screening typically structures these models as dynamic games—for example, a strategic interaction between one worker and n firms, complete with an order of play—and solves these models by characterizing strategies (and beliefs) falling under certain equilibrium notions. Traditionally, the informed party (e.g. the worker) will typically move first (e.g. by getting education) in *signaling* games and move second in *screening* games. The first is typically a dynamic Bayesian game—the first player taking an *action* and the second player forming a *belief*, updating it upon seeing the action, and responding with their own—with unrestricted beliefs leading to infinitely many PBE. In screening, the uninformed players—the n firms—act first by offering a menu of contracts, followed by the informed player choosing those contracts. When considering competitive markets, however, it often makes little sense to stipulate a strategic interaction with one single informed player, and with one side necessarily moving first, so while selection market theory appropriates game theoretic equilibrium concepts, it does so with little emphasis or fidelity to the game theoretic details.¹

1. For example, we just as commonly say that “firms *screen* workers” (rather than “workers signal to firms”). When a worker chooses a level of education, a menu of contracts is already being offered—as represented through the price of labour associated with each level of education. These prices essentially represent the beliefs of the market, but unlike signaling games where the set of

A seminal work on competitive markets with asymmetric information is Riley (1979), in which n competitive uninformed parties offer a market menu of contracts. The *reactive equilibrium* concept, introduced in the paper, consists of a menu which, if any uninformed party were to undercut in price, would induce a *reaction* which renders the undercutting unprofitable. Riley’s model is similar to Rothschild-Stiglitz in that any pool can be separated by cream-skimmers, but different in that the Pareto-dominant zero-profit separating equilibrium² cannot be de-stabilized by any pooling or separating market offering without that offering being threatened with unprofitability by a market reaction. Subsequent work in game theoretic (Cho and Kreps 1987; Banks and Sobel 1987) and price theoretic (Azevedo and Gottlieb 2017) models of behaviour under asymmetric information also suggested the same outcome, now known as the *Riley outcome*.

1.2 Recent Advances and Challenges

For a number of decades after the 1970s, Akerlof (1970), Spence (1973), and Rothschild and Stiglitz (1976) stood as the workhorse models in empirical work and benchmark models in theoretical work. In a handbook chapter on connecting theory to data in social insurance, Chetty and Finkelstein (2013) undertake the theoretical discussion almost entirely using an Akerlof (1970) model, only augmenting it as needed. And in a Yale Law Journal article arguing that adverse selection is an exaggerated threat, Siegelman (2003) presents the theory in the Rothschild and Stiglitz (1976) paradigm, then proceeds to provide empirical evidence against it and point to weaknesses in the theory.

Siegelman was not alone: an empirical strand of literature which emerged in the 2000’s, beginning with the pioneering work of Chiappori and Salanié (2000), sought

admissible beliefs may be large or infinite (depending on how one defines an equilibrium, more on this later), prices in the competitive market model are set mechanically by a zero profit condition.

2. or ‘*least-cost* separating equilibrium’

to test adverse selection theory looking for the correlation between risk and coverage which the theory predicts—the so-called “positive correlation test.” What was surprising was how often the test produced a null result. Subsequently, Finkelstein and McGarry (2006) pointed out that null results could be explained by multiple dimensions of private information. Cutler, Finkelstein, and McGarry (2008) investigate the issue further, showing the importance not only of private information about risk, but the second dimension of risk *tolerance*. In particular, while it is true that *ceteris paribus*, high-risk people select into high coverage, it is also true that risk tolerant people select into low coverage. Because risk tolerant people tended to be high risk, the competing forces of adverse and advantageous selection led to the lack of positive correlation between coverage and risk.

This highlighted the importance of higher dimensions in selection markets—an aspect conspicuously missing from selection theory, which primarily focused on one dimension of contracts (e.g. coverage) and one dimension of types (e.g. risk), both ordered.³ It was in this context that the Veiga and Weyl (2016) and Azevedo and Gottlieb (2017) models were developed for selection markets with multidimensional heterogeneity.

A second rationale for a lack of adverse selection observed in markets is the possibility that, ironically, selection markets are selective, in that the more adversely selected a market is, the more likely it is to not exist (and thereby cannot be observed) because, as we show in chapter 3, such markets take only small frictions to collapse. Some of the empirical literature in the last ten years has been dedicated to studying selection in markets which do not exist. One notable avenue, pursued by Hendren (2013, 2017), has been to identify a theoretical “no-trade” condition under which a market would be so adversely selected that it could not exist, and then test empirically whether the condition is met; if so, then private information precludes the market from existing.

3. See Azevedo and Gottlieb (2019) for a non-existence result for reactive equilibria when types are not ordered along one dimension.

Looking forward, one final incongruence between theory and empirics which has developed in recent years has to do with the fact that all the above-described models assume perfect rationality, while much of the literature of the last decade—particularly in health insurance—has pointed to the empirical relevance of choice frictions (Chandra, Handel, and Schwartzstein 2019). Just as the empirical relevance of multidimensional heterogeneity presented the need for new theoretical tools, so too does the empirical relevance of choice frictions.

1.3 Overview of This Thesis

Chapter 2 provides a basic result about the equilibrium of Azevedo and Gottlieb (2017) which reconciles it with the Riley outcome. Typically, an easy and general way to solve for the outcome is via an initial value problem (IVP) first presented in Riley (1979) (for a recent example of its use, see Chen, Ishida, and Suen (2022)). I derive the conditions under which it also produces a unique Azevedo-Gottlieb outcome. The mathematical contribution here is two-fold. First, I connect a mathematical tool commonly used in economics—the differential equation representing an incentive-compatible separating equilibrium—with a modern equilibrium concept. Second, I show that general propositions about the Azevedo-Gottlieb equilibrium concept can be proven without using the definition of the equilibrium. The concept has always been vulnerable to the criticism that it is mathematically intractable; the equilibrium has a complicated definition, and because the paper provides only theorems regarding existence and necessary conditions, empiricists and theoreticians alike may forgo this concept because it is not clear that one can prove a given proposition without starting from the definition. Even if the proposition was rather intuitive, one risks being bogged down in the details. Indeed, the paper itself proves only specific examples with the given theorems. On the other hand, I prove a more general proposition using only the theorems provided in the paper.

Chapter 3 studies a Riley model of insurance coverage with a fixed cost of provid-

ing contracts, solves for the D1 equilibrium (Banks and Sobel 1987), and reconciles it with a necessary condition for no-trade in Hendren (2013) and Azevedo and Gottlieb (2017). I show the equilibrium is given by a pool at zero coverage, and possibly a separating portion, defined by a Riley IVP. The parameters of the specified model determine an open set in contract-price space, and the price function of the separating contracts is determined by the IVP over its maximal interval of existence in that open set. Of course, if the open set is $\emptyset$, then there is only a pool. This yields a somewhat intuitive condition for no trade: if and only if the riskiest type is not willing to pay the fair price of full coverage.

Chapter 4 presents a framework for simulating selection markets with multiple dimensions of information and of contracts which allows for the incorporation of choice frictions. I then construct a model within the framework to study the implications of rational inattention on selection markets. The contribution of this chapter spans several fronts: i) I construct a new approach to modeling selection markets—a dynamic multi-dimensional model with stochastic choices—which contrasts with and supplements the standard static (and often uni-dimensional) models with fully rational agents ii) I provide a new, intuitive, market-based foundation for D1/Riley equilibria and iii) I find that outcomes in markets with inattentive agents can diverge significantly from markets with perfectly rational agents.

Chapter 2

Reconciling Azevedo and Gottlieb’s Equilibrium with Riley’s Initial Value Problem

2.1 Introduction

When considering economic problems with asymmetric information—ranging from job market signaling to adverse selection in insurance markets to strategic communication—it is often useful to first consider a simple model with an interval of types and an interval of *contracts*.⁴ Given *single-crossing* preferences (Spence 1973; Mirrlees 1971), the initial value problem (IVP) of Riley (1979) describes a Pareto-dominant (or “least cost”) separating equilibrium where i) sellers make zero profits and ii) buyers set marginal cost to equal marginal rate of substitution between money and the signal given by the contract choice.

Recently, interest has shifted towards the equilibrium concept of Azevedo and Gottlieb (2017, hereafter, AG). This equilibrium is attractive in that it purportedly

4. Depending on context, *contracts* might be replaced with *signals*, *messages*, *actions*, etc. as the object of interest

replicates the Pareto-dominant separating equilibrium when both exist, but unlike the latter, the former is capable of handling multiple dimensions and unordered types. However, to my knowledge, no theoretical results have been proven which reconciles generally the AG model with Riley's IVP, and so before we can write down the IVP and assume that it yields the AG equilibrium, several questions must be answered. Does the IVP have a unique solution under the softer assumptions of AG? Does the solution yield a valid or unique equilibrium? Are there non-separating equilibria which the IVP fails to characterize? What about semi-separating or mixed-strategy equilibria?

This chapter addresses these questions, and draws some familiar conclusions. With a local Lipschitz condition on the marginal rates of substitution, equilibria are (in some sense) unique and almost every type is allocated to a single contract. When the contract space is sufficiently large, as is sometimes explicitly assumed, types separate; otherwise, bad types separate and good types pool, as in Cho and Sobel (1990) or Kartik (2009). Semi-separating/pooling equilibria depend on the distribution of types, while the separating equilibrium depends only on the support. Separating, semi-separating, and pooling equilibria are (almost) disjoint occurrences. Thus, it suffices to solve the IVP first and check the contract space is large enough after the fact. Without the Lipschitz condition, however, IVPs are not unique and an AG equilibrium corresponding to one need not be Pareto-dominant.

The contribution of this chapter is two-fold. First, I reconcile AG with the Riley IVP. Second, in doing so, I show how certain results may be obtained using only the theorems provided in AG, which include only an existence condition and some necessary but not sufficient conditions. One concern among economists has been that without sufficient conditions, one may have to work with the incredibly complex machinery of AG in order to prove any theorems about the equilibrium. To my knowledge, the theorems have not been used to prove any general results. Indeed, AG themselves only use their theorems to characterize specific examples.

Relation to the Literature. The seminal paper by Riley (1979) introduced the concept of a *reactive* equilibrium; in lay terms, an equilibrium is reactive if for every strictly profitable deviation, there exists a profitable market *reaction* which renders the deviation strictly unprofitable. Under stronger assumptions, Riley also provided the IVP presented herein for the Pareto-dominant separating equilibrium, which coincided with the reactive equilibrium. Engers and Fernandez (1987, hereafter, EF) generalised the Riley model and relaxed some assumptions; but left the assumption that the contract space was so large so as to guarantee a separating equilibrium. Moreover, whether the equilibrium satisfied an IVP was outside the scope of the paper.

Up until now, a workable approach has been to restrict attention to “sufficiently large” compact contract spaces with a C^2 utility function U . One then appeals to the Mailath (1987) result that any incentive-compatible (IC) separating allocation must satisfy an ordinary differential equation (ODE). However, $U \in C^2$ does not guarantee uniqueness of the solution to the IVP, nor necessarily imply the non-existence of IC semi-separating states, and so one must then refine away all other equilibrium—typically, leaving behind the Pareto-dominant separating one. Moreover, it is possible, depending on the marginal rates of substitution, that the allocations “explode” to infinity too quickly, thus even an unbounded action space can fail to host a separating equilibrium.

This chapter relaxes some of these assumptions and characterizes the equilibria. While AG is both technically unwieldy and does not explicitly prove that their equilibrium aligns with the traditionally-expected outcomes of single-crossing models—say, Riley or D1 equilibrium (Banks and Sobel 1987)—it does provide very useful necessary conditions and an existence guarantee, which is sufficient to parlay into a uniqueness result. As we will see, what guarantees uniqueness is not $U \in C^2$, but the local Lipschitz property of the marginal rate of substitution $-U_x/U_p$.

2.2 Model

Continuum screening environments. A screening environment consists of a type space $\Theta = [0, \bar{\theta}]$, a type distribution F with support Θ on a set probability space $(\Omega, \mathcal{F}, P)$, a contract (or ‘signal’) space $X = [0, \bar{x}]$, a utility function $U : X' \times R' \times \Theta' \rightarrow \mathbb{R}$, and cost function $c : X' \times \Theta' \rightarrow \mathbb{R}_+$, where $X' \supset X$, $R' \supset \mathbb{R}_+$, and $\Theta' \supset \Theta$ are each open sets.⁵

Assumption 1. $U, c \in C^1$.

Assumption 2. Without loss of generality,⁶ $U_x < 0$, $U_p < 0$, $c_x \leq 0$, $c_\theta < 0$, and $0 \geq c_x(x, \theta) > -\frac{U_x(x, c(x, \theta), \theta)}{U_p(x, c(x, \theta), \theta)} > -\infty$.⁷

Assumption 3. Indifference curves are single-crossing and those of better types are less steep; that is:

$$\theta > \theta' \implies -\frac{U_x(x, p, \theta)}{U_p(x, p, \theta)} > -\frac{U_x(x, p, \theta')}{U_p(x, p, \theta')}.$$

Assumption 4. There exists a metric d on X and a number L for which X is compact and $U(x, p, \theta) \geq U(x', p', \theta), p' > p \implies p' - p \leq L \cdot d(x, x')$. That is to say that the marginal rates of substitution are bounded.⁸

Equilibrium notions. A *mechanism* (x, p) consists of an allocation $x : \Theta \rightarrow X$ and a price $p : X \rightarrow \mathbb{R}_+$. The mechanism is incentive compatible (IC) if

$$x(\theta) \in \arg \max_{y \in X} U(y, p(y), \theta).$$

5. We situate the problem in open sets merely for differential equations purposes. Another approach, taken by Mailath and Von Thadden (2013), is to use the one-sided derivative where appropriate.

6. i.e. these inequalities can be appropriately reversed

7. This is to say that signaling is not first-best. In the context of insurance where $1 - x$ is coverage and $1 - \theta$ is probability of loss, $c_x > U_x/U_p$ if insurers are risk-neutral, insureds are risk-averse, and $0 < \theta < 1$. In the context of job market signaling where education is x and productivity is $-c$, $c_x > U_x/U_p$ if the productivity gain is not sufficient to justify additional education.

8. Except when $\bar{x} = \infty$, d can be taken to be the usual metric $|x - y|$.

Notice we deal only with pure-strategy allocations. An *equilibrium* is an IC mechanism. y is a *pool* of x if $x^{-1}(y)$ has more than one element (x allocates more than one type to y). The equilibrium is *separating* if x has no pools (x is one-to-one, allocating each type to a unique element), *pooling* if it has one pool ($x(\Theta)$ is a singleton, x maps all types into a single element), and *semi-separating* otherwise (x maps types to two or more elements, but at least two types are allocated to the same element). The equilibrium is *zero-profit* if, for any contract y in the image of the allocation $x(\Theta)$, $E[c(y, \theta) | \theta \in x^{-1}(y)] = p(y)$.⁹

In this chapter, we will consider (strong) Azevedo-Gottlieb equilibria (SAGE). The reader may refer to AG for a technical exposition of SAGE, but for the purposes of this chapter, it suffices to know the following. SAGE is zero-profit and IC (AG, Proposition 1.1), with a Lipschitz continuous price function (AG, Proposition 1.3), such that (to quote AG, Proposition 1.2 directly) “every contract that is not traded in equilibrium has a low enough price for some consumer to be indifferent between buying it or not, and the cost of this consumer is at least as high as the price.”

2.3 On the Generality of the Model

Here, I prove two facts about the model. First, it suffices to restrict attention to non-probabilistic allocations $x : \Theta \rightarrow X$. See, formally, AG allocations consist of a measure $\alpha : \mathcal{B}_X \otimes \mathcal{B}_\Theta \rightarrow \mathbb{R}$ on a product space. To bridge the two, and to allow the latter to take advantage of the theory of the former, we use lemma 1.

Second, the model generalizes the traditional signaling setup where players are evaluated on their expected type, rather than their expected cost given their contract. Another viewpoint is that the zero-profit equilibrium is essentially a productive signaling equilibrium, as shown in lemma 2. In productive signaling, the payoff is

9. Note this is *zero profit in expectation*, not zero profit for each contract. In competitive markets, all equilibria will be zero-profit, although not screening environments in general in the literature. See monopolist pricing models, for example.

$U(x, E_{t \sim \mu(x, \cdot)}[c(x, t)], \theta)$ where $\mu(x, \cdot)$ is a regular conditional probability representing the belief (a distribution over θ) of the signal receiver upon observing signal x (in AG, μ would be given by α in equilibrium) and U is the signal sender's utility. This generalizes unproductive signaling Perfect Bayesian Equilibria, where the payoff is $V(x, E_{t \sim \mu(x, \cdot)}[t], \theta)$, a utility function given by expected type (often called "reputation") rather than expected cost.

Lemma 1. *In the SAGE of this model, types are allocated to one non-pool **at most** (in an a.e. sense with respect to α). Moreover, if there is a pool, it must be $\bar{x}$. Lastly, only one type can be allocated to both a non-pool and $\bar{x}$.*

Proof. First, assume for contradiction that θ is allocated to two non-pools $x_1 < x_2$. Since they are not pools, it must be the case that $p(x_1) = c(x_1, \theta)$ and $p(x_2) = c(x_2, \theta)$. To see that θ *strictly* prefers x_1 , observe that

$$\begin{aligned} U(x_1, c(x_1, \theta), \theta) - U(x_2, c(x_2, \theta), \theta) &= \int_{x_1}^{x_2} \frac{dU}{dx}(x, c(x, \theta), \theta) dx \\ &= \int_{x_1}^{x_2} U_x + U_p c_x dx > 0 \end{aligned}$$

by the assumption that $c_x > -U_x/U_p$ and $U_p < 0$. Hence, IC is violated.

Next, assume for contradiction that there is a pool $y < \bar{x}$. It follows there are types t in the pool whose cost $c(y, t)$ is strictly less than $p(y) = E[c(y, \theta) | \alpha]$. Let t_* be the infimum of such types and let $q_*(x)$ be the indifference curve of t_* through $(y, p(y))$. Since p is continuous (AG 1.3), we can pick $\varepsilon > 0$ such that $p(y + \varepsilon) > c(y + \varepsilon, t_*)$. IC implies that $p(y + \varepsilon) \geq q_*(y + \varepsilon)$. In light of AG 1.2, suppose θ that some type is indifferent between buying $y + \varepsilon$ or not. Given that $p(y + \varepsilon) \geq q_*(y + \varepsilon)$, by single-crossing, $\theta \geq t^*$. But then, $c(y + \varepsilon, \theta) < p(y + \varepsilon)$, contradicting AG 1.2.

Lastly, suppose both $\theta' < \theta''$ are allocated to the pool $\bar{x}$ and the non-pools y' and y'' respectively. By IC, θ' must be indifferent between $(\bar{x}, p(\bar{x}))$ and $(y', p(y'))$,

which by single-crossing, implies θ'' strictly prefers $(y', p(y'))$, violating IC. $\square$

Lemma 2 (Representation lemma). *Let $V(x, t, \theta)$ be differentiable and strictly monotonic in x and reputation t . Without loss of generality assume $V_x < 0$ and $V_t > 0$. Then there exists a set $\mathcal{P}$ and functions $U : X \times \mathcal{P} \times \Theta \rightarrow \mathbb{R}$ and $c : X \times \Theta \rightarrow \mathcal{P}$, such that $U_x < 0$, $U_p < 0$, $c_x \leq 0$, $c_\theta < 0$, and $V(x, t, \theta) = U(x, c(x, t), \theta)$. As a consequence, any allocation-belief pair (x, μ) satisfying*

1. μ is a regular conditional distribution (RCD) of θ given $x \circ \theta$; that is, μ satisfies Bayes' rule

2. $x(\theta) \in \arg \max_{y \in X} V(y, E_{t \sim \mu(y, \cdot)}[t], \theta)$ for all $\theta \in \Theta$; that is, x is IC

also satisfies the condition $x(\theta) \in \arg \max_{y \in X} U(y, E[c(y, t)|t \in x^{-1}(y)], \theta)$ for all $\theta \in \Theta$.

Proof. Fix V . Set $c(x, \theta) = -\theta$ and $U(x, p, \theta) = V(x, -p, \theta)$. It follows $U_x = V_x < 0$, $U_p = -V_t < 0$, $c_x = 0 \leq 0$, $c_\theta = -1 < 0$, and $U(x, c(x, t), \theta) = U(x, -t, \theta) = V(x, t, \theta)$. $E_{t \sim \mu(y, \cdot)}[t] = E[c(y, t)|t \in x^{-1}(y)]$ holds by construction of the RCD (Kallenberg 2002, Proposition 8.5). $\square$

2.4 Results

Definition 1. The ordinary differential equation (ODE) for allocations is

$$\xi'(\theta) = -\frac{U_p c_\theta}{U_x + U_p c_x} =: g(\theta, \xi(\theta)) \quad (\text{ODE:A})$$

evaluated at $x = \xi(\theta)$, $p = c(\xi(\theta), \theta)$, and $\theta = \theta$.¹⁰ The IVP for allocations is given by (ODE:A) and the initial condition $\xi(0) = 0$.

Definition 2. The ODE for price functions is

$$\phi'(x) = -\frac{U_x(x, \phi(x), \theta(x, \phi(x)))}{U_p(x, \phi(x), \theta(x, \phi(x)))} =: f(x, \phi(x)) \quad (\text{ODE:P})$$

10. that is, $\xi'(\theta) = -\frac{U_p(x(\theta), c(\xi(\theta), \theta), \theta) c_\theta(\xi(\theta), \theta)}{U_x(\xi(\theta), c(\xi(\theta), \theta), \theta) + U_p(\xi(\theta), c(\xi(\theta), \theta), \theta) c_x(\xi(\theta), \theta)}$

where $\theta(\cdot, \cdot)$ is defined implicitly by $c(x, \theta(x, p)) = p$. The IVP for price functions is given by (ODE:P) and the initial condition $\phi(0) = c(0, 0)$.

The *implicit* ODE for price functions is given by the chain rule on $\phi(x(\theta)) = c(x(\theta), \theta)$:

$$\phi'(x) = c_x + c_\theta/x' = c_x - \frac{(U_x + U_p c_x)c_\theta}{U_p c_\theta} = -\frac{U_x}{U_p}$$

Definition 3. The SAGE is *unique* if the price function is unique.

With definitions in tow, I now proceed to the main results. The proofs of the results are deferred to Section 2.6. I will first present an overview of the intuition. If we view x as a signal, the worst type begins by sending the weakest signal. Each type θ sends a signal just large enough to “separate” from a marginally inferior type $\theta - \partial\theta$ (conversely, one could imagine that a cream skimmer would lure θ away from $\theta - \partial\theta$ with a low enough price for a marginally larger signal). At each turn, θ asks whether they would prefer to separate and pay their actuarially fair price, or pool at $\bar{x}$ and pay a price $E[c(\bar{x}, t)|t > \theta]$. If θ^* is indifferent between the two, all types $\theta > \theta^*$ strictly prefer to pool, and thus pool.

If the contract space has enough “room” for types to separate—that is, $\bar{x}$ is sufficiently large that no one wants to pool—then pooling does not occur. When the space is sufficiently restricted, high types begin to pool at $\bar{x}$. When the space is further restricted, all types pool at $\bar{x}$.

A key insight used to show the uniqueness of SAGE is that because types choose contracts by setting their indifference curve tangent to the graph of p , p must be absolutely continuous, otherwise p' would equal $-\infty$ in some places. AG, Proposition 1.4 guarantees that p is a.e. differentiable, and so jointly, this guarantees that p satisfies (ODE:P). Local Lipschitz continuity guarantees the uniqueness of solutions for the ODE.

Proposition 3. *If the slope field f of (ODE:P) is locally Lipschitz in p , then the following statements are equivalent:*

1. *There is a separating SAGE.*
2. *There is a unique separating SAGE (x, p) where p solves the IVP for price functions on $x(\Theta)$.*
3. *Let ϕ solve the IVP for price functions up to its maximal interval of existence I . There is an $x \in I$ where $x \leq \bar{x}$ for which $\phi(x) = c(x, \bar{\theta})$.*
4. *There are no semi-separating/pooling SAGE.*

Proposition 4. *Suppose f is locally Lipschitz in p , any of the statements in proposition 3 are violated, and $U(\bar{x}, E[c(\bar{x}, \theta)], 0) > U(0, c(0, 0), 0)$. The SAGE consists of a pool at $\bar{x}$.*

The semi-separating SAGE can be constructed as follows. By the Implicit Function Theorems for differentiable and for strictly monotonic functions (Hazewinkel et al. 1990), for each $\theta \in \Theta$, we can construct $q(x; \theta)$, the indifference curve of θ passing through the pool, implicitly via

$$U(x, q(x; \theta), \theta) - U(\bar{x}, E[c(\bar{x}, t)|t > \theta], \theta) = 0$$

and construct $\xi(\theta)$, the separating contract for which θ would be indifferent to the pool, via

$$U(\xi(\theta), c(\xi(\theta), \theta), \theta) - U(\bar{x}, E[c(\bar{x}, t)|t > \theta], \theta) = 0.$$

Further, ξ is continuous and $q_x(x; \theta) = -\frac{U_x(x, q(x; \theta), \theta)}{U_p(x, q(x; \theta), \theta)}$; q is also continuous in θ because

$$\begin{aligned} \lim_{\theta \rightarrow \theta^-} U(x, q(x; \theta), \theta) &= \lim_{\theta \rightarrow \theta^-} U(\bar{x}, E[c(\bar{x}, t)|t > \theta], \theta) \\ &= \lim_{\theta \rightarrow \theta^+} U(\bar{x}, E[c(\bar{x}, t)|t > \theta], \theta) \\ &= \lim_{\theta \rightarrow \theta^+} U(x, q(x; \theta), \theta). \end{aligned}$$

Define $D = \{(x, p) : \exists \theta \in \Theta \text{ s.t. } x \geq \xi(\theta) \text{ and } p \geq c(\xi(\theta), \theta)\}$ to be the region of $X' \times R'$ (see the model assumptions) northeast of $L := \{(\xi(\theta), c(\xi(\theta), \theta)) : \theta \in \Theta\}$. Define $t : D \rightarrow \Theta$ implicitly by $q(x, t(x, p)) - p = 0$, which we can do since for fixed x , $q(x, \theta) - p$ is strictly monotonic in p , and vice versa. Finally, define the slope field $h : X' \times R' \rightarrow \mathbb{R}$ by

$$h(x, p) = \begin{cases} f(x, p) & (x, p) \in D^c, \\ q_x(x, t(x, p)) & (x, p) \in D. \end{cases} \quad (2.1)$$

Proposition 5. *Suppose f is locally Lipschitz in p , any of the statements in proposition 3 are violated, and $U(\bar{x}, E[c(\bar{x}, \theta)], 0) < U(0, c(0, 0), 0)$. If further, F is non-atomic, there is a unique SAGE price function satisfying $p'(x) = h(x, p)$, $p(0) = c(0, 0)$.*

That F is non-atomic is sufficient for SAGE to not need probabilistic allocations, but if $\hat{\theta}$ has a positive mass, it is possible that there is no SAGE allocation where $\hat{\theta}$ is allocated a.s. to one contract.

The final proposition deals with whether AG and EF select for a unique and, in some sense, ‘logical’ outcome when f is not locally Lipschitz. It hints at something interesting: EF equilibrium loses its existence guarantees when types are unordered, as shown by Azevedo and Gottlieb (2019), but on the other hand, AG might not select for a desirable equilibrium even when types are ordered, whereas EF does.

Proposition 6. *There is always a solution to the IVP in the sense of Carathéodory. Let $\mathcal{F}$ be the set of solutions. $\bar{\phi}(x) := \sup_{\phi \in \mathcal{F}} \phi(x)$ and $\underline{\phi}(x) := \inf_{\phi \in \mathcal{F}} \phi(x)$ are both in $\mathcal{F}$. SAGE sometimes selects for $\bar{p}$, whereas if the assumptions in EF are satisfied, then EF always selects the Pareto-dominant market offering.*

Note that this is purely an artifact of the model. AG tends to select for separating equilibria with a least-cost initial value. Any price function solving the IVP, such as $\bar{\phi}$, must be associated with such an equilibrium where the price function equals $\bar{\phi}$ on-the-equilibrium, since the only off-the-equilibrium contracts are $(x(\bar{\theta}), \bar{x}]$. This does

not happen with discrete type spaces: tracing along indifference curves to identify the least-cost separating equilibrium yields a unique equilibrium.

2.5 Discussion and Concluding Remarks

I have thus far postponed the discussion of which assumptions are strictly necessary and which are merely simplifying. Purely for the purposes of ODEs, we extend the definition of U and c to open sets X' , R' , and Θ' . The assumption that types are lower bounded is to assume there exists a worst type. Likewise, there is a ‘worst’ contract in terms of signaling value which is ‘best’ in terms of consumer utility. It is sufficient for signaling purposes that worst types and contracts *exist*, but it is arbitrary that lower types are worse types. Without loss of generality, we assume the lower bounds on X and Θ to be zero. The contract space may be the extended nonnegative reals $[0, \infty]$, in which case we include the element ∞ .¹¹ The choice made that $c_x \leq 0$ and $c_\theta < 0$ is arbitrary; in the Spence (1973) job market signaling game with a productivity function $v(x, \theta)$ such that $v_x \geq 0$ and $v_\theta > 0$, one could set $c = -v$. $c \geq 0$ is necessary for SAGE, but it not restrictive because even if c is not lower bounded, we may monotonically transform the range of c and accordingly transform U . The assumption that $c_x > -U_x/U_p$ implies that signaling is *costly*, even if somewhat productive: recall that, because zero profits are made, productivity gains from higher x are returned in the form of reduced c , with the marginal cost reduction being c_x , and that must not fully offset the utility loss $-U_x/U_p$. The marginal rate of substitution being Lipschitz along the isocost curves is used to invoke the Picard-Lindelöf Theorem for uniqueness of IVP solutions.

One notable motif of these types of models is that they have no individual rationality (IR) constraint: the rationale is that, provided the contract space is sufficiently large, the equilibrium is dependent only on the *support* of the distribution, not the distribution itself, and therefore, provided IR constraints do not change the support

11. under the standard assumptions concerning $\infty \in \overline{\mathbb{R}}$; see, for example, Folland (1999)

of the distribution, the equilibrium is otherwise unaffected. This is not always the case for semi-separating or pooling equilibria.

Nevertheless, the main idea of this chapter is that the economist may work out the IVP for price functions first (solving directly or numerically), and check that the solution satisfies condition 3 in proposition 3 *ex post facto*. If so, then the economist can rest assured that the equilibrium is separating, the solution *is* the price function of the informational equilibrium, and the equilibrium is unique.

2.6 Chapter Appendix: Omitted Proofs

Lemma 7. *Let $I \subset \Theta$ be an open interval and x be a SAGE allocation. If x is one-to-one on I , then p is strictly decreasing, x is continuous and therefore strictly increasing on I .*

Proof. First, if $x(\theta') > x(\theta'')$ and $p(x(\theta')) \geq p(x(\theta''))$, then θ' would prefer $x(\theta'')$.

Thus, if x has a discontinuity at θ , it must be an increasing one. Let $t_n \downarrow \theta$ but $\lim_{n \rightarrow \infty} x(t_n) > \lim_{t \rightarrow \theta^-} x(t) \equiv x(\theta-)$. Then:

$$\begin{aligned}
& \lim_{n \rightarrow \infty} U(x(t_n), c(x(t_n), t_n), t_n) - U(x(\theta-), c(x(\theta-), \theta-), t_n) \\
&= \lim_{n \rightarrow \infty} U(x(t_n), c(x(t_n), t_n), t_n) - U(x(\theta-), c(x(\theta-), \textcolor{red}{t}_n), t_n) \\
&= \lim_{n \rightarrow \infty} \int_{x(\theta-)}^{x(t_n)} \frac{dU}{dx}(x, c(x, t_n), t_n) dx \\
&\leq \lim_{n \rightarrow \infty} \int_{x(\theta-)}^{\textcolor{red}{x}(\theta)+\varepsilon} \frac{dU}{dx}(x, c(x, t_n), t_n) dx \\
&\leq \int_{x(\theta-)}^{x(\theta-)+\varepsilon} \left[\limsup_{n \rightarrow \infty} \frac{dU}{dx}(x, c(x, t_n), t_n) \right] dx \\
&= \int_{x(\theta-)}^{x(\theta-)+\varepsilon} [U_x(x, c(x, \theta), \theta) + U_p(x, c(x, \theta), \theta) \cdot c_x(x, \theta)] dx < 0,
\end{aligned}$$

which implies for some $(x(\theta-), p(x(\theta-))) \succ_{t_n} (x(t_n), p(x(t_n)))$ for some t_n , contradicting IC. Continuous and one-to-one implies strictly monotonic, and single-crossing implies x is not strictly decreasing. $\square$

Lemma 8. *Let (x, p) be a SAGE. If $I \subset \Theta$ is an interval, and x is one-to-one on I , then p is absolutely continuous on $x(I)$. In particular, if x is separating on Θ , p is absolutely continuous.*

Proof. Assume for contradiction that p is not absolutely continuous. By AG, Proposition 1.3 it is nevertheless continuous. Moreover, because $p(x(\theta)) = c(x(\theta), \theta)$, p must be strictly decreasing on $\overline{x(I)}$. x is continuous, and I is an interval and therefore connected, so $\overline{x(I)}$ is also an interval.

It follows that the set $\{x \in \overline{x(I)} : p'(x) = -\infty\}$ has positive Lebesgue measure (Leoni 2017, 3.46). Pick a $\theta^* \in I$ for which $x(\theta^*) \in \text{int}(x(I))$ and $p'(x(\theta^*)) = -\infty$, and it follows from the fact that the marginal rate of substitution is bounded that the indifference curve of θ^* crosses the graph of p , and therefore, the mechanism is not IC, contradicting AG, Proposition 1.1. $\square$

Lemma 9. *Let (x, p) be a SAGE. If $x(0)$ is not in a pool, then $x(0) = 0$, and therefore $p(0) = c(0, 0)$.*

Proof. $x(0) > 0$ and $p(x(0)) = c(x(0), 0)$ would mean every $x \in [0, x(0))$ contradicts AG, Proposition 1.2. $\square$

Lemma 10. *Any pool in SAGE consists of an interval of types.*

Proof. Let $(y, p(y))$ be a pool and let $\theta < \theta' < \theta''$. Let $S_y = \{(x, p) : (x, p) \succ_{\theta} (y, p(y))\}$ and define S'_y and S''_y likewise. By single-crossing, $\overline{S'_y} \cap (S_y \cup S''_y)^c = \{(y, p(y))\}$. *Id est*, the only pair weakly preferred over $(y, p(y))$ by θ' but not strongly preferred by θ or θ'' is $(y, p(y))$ itself. $\square$

Corollary 10.1. *Any semi-separating SAGE consists of a pooling interval containing $\bar{\theta}$ and a separating interval containing 0.*

Lemma 11. *If $f(x) = g(x)$, $f'(x) < g'(x)$, then for some $\delta > 0$, $|h| < \delta$ implies $f(x - h) < g(x - h)$ and $f(x + h) > g(x + h)$. That is, f and g “cross.”*

Proof. Set $\phi = g - f$. For $\varepsilon = \phi'(x)/2 > 0$, we can take $\delta > 0$ such that $0 < |h| < \delta \implies |\phi'(x) - \frac{\phi(x+h) - \phi(x)}{h}| < \phi'(x)/2$, implying $\phi'(x)/2 < \phi(x+h)/h < 3\phi'(x)/2$. $\square$

Proof of proposition 3. (1 $\implies$ 2) Let (x, p) be a separating SAGE. By lemma 8, p is absolutely continuous. By AG, Proposition 1.4, p is a.e. differentiable. Where it is differentiable on $x(\Theta)$, it must satisfy the ODE, as otherwise, the indifference curve of some θ would cross the graph of p at $x(\theta)$, per lemma 11. By lemma 9, the initial condition is also unique: $p(0) = c(0, 0)$. So p must solve the IVP uniquely on $x(\Theta)$ (Hale 1980). For non-traded contracts $x > x(\bar{\theta})$, AG, Proposition 1.2 implies the price function traces along the indifference curve of $\bar{\theta}$ until 0 is met, after which $p = 0$. To see that x must also be unique, consider two possible allocations x_1 and x_2 ; recall that $p' = -U_x/U_p < c_X$ implies p crosses each isocost curve $c(\cdot, \theta)$ only once, so for each θ , there is only one x for which $p(x) = c(x, \theta)$, thus $x_1 = x_2$.

(1 $\implies$ 3) Let (x, p) be a SAGE and let p, ϕ solve the IVP. By Picard-Lindelöf, $p = \phi$, so $\phi(x(\bar{\theta})) = p(x(\bar{\theta})) = c(x(\bar{\theta}), \bar{\theta})$.

(3 $\implies$ 2) Let $\phi(x) = c(x, \bar{\theta})$. If (x, p) is a semi-separating SAGE, then there is a separating interval $I = [0, \theta^*)$ and a pooling interval $(\theta^*, \bar{\theta}]$. By lemma 8, p is absolutely continuous on $x(I)$; as before, combined with lemma 9 this implies $p = \phi$ on $x(I)$. On the pooling interval, $x(\theta) = \bar{x}$ and $p(\bar{x}) = E[c(\bar{x}, t) | t > \theta^*]$, such that θ^* is indifferent between $\lim_{\theta \rightarrow \theta^*} x(\theta)$ and $\bar{x}$. But by construction of ϕ ,

$$U(x(\theta^*), \phi(x(\theta^*)), \theta^*) \geq U(x, \phi(x), \theta^*) = U(x, c(x, \bar{\theta}), \theta^*).$$

Recall that $c_x > -U_x/U_p$ implies $U(x, c(x, \bar{\theta}), \theta^*) \geq U(\bar{x}, c(\bar{x}, \bar{\theta}), \theta^*)$, which, by monotonicity of U , is $> U(\bar{x}, E[c(\bar{x}, t) | t > \theta^*], \theta^*)$. Ergo, θ^* strictly prefers $\lim_{\theta \rightarrow \theta^*} x(\theta)$ to $\bar{x}$. A nearly identical proof applies to show that no pooling SAGE exists either.

(4 $\implies$ 1) SAGE always exists, so if it isn't semi-separating or pooling, it must be separating. $\square$

Proof of proposition 4. From lemmas 1 and 10 and proposition 3, we know the SAGE is not separating, and is not semi-separating if $x(0)$ is in the pool at $\bar{x}$. Assume for contradiction that $x(0)$ is not in the pool at $\bar{x}$. Lemma 9 implies $x(0) = 0$ and $p(x(0)) = c(0, 0)$. Clearly, $E[c(\bar{x}, \theta) | \theta > t]$ is strictly decreasing in t since $c_\theta < 0$; therefore,

$$U(\bar{x}, E[c(\bar{x}, \theta) | \theta > \theta^*], 0) > U(\bar{x}, E[c(\bar{x}, \theta)], 0) > U(0, c(0, 0), 0)$$

where θ^* separates the separating and pooling intervals, contradicting IC. Therefore, $x(0)$ is in the pool at $\bar{x}$, and the result follows. $\square$

Lemma 12. *Let θ^* divide the pooling and separating intervals of a semi-separating SAGE. Then*

$$\lim_{\theta \rightarrow \theta^* -} U(x(\theta), c(x(\theta), \theta), \theta^*) = U(\bar{x}, E[c(\bar{x}, \theta) | \theta > \theta^*], \theta^*)$$

Proof. If θ^* strictly preferred $\lim_{\theta \rightarrow \theta^* -} x(\theta)$ to $\bar{x}$, then by continuity, so would some type $\theta^* + \varepsilon$, contradicting IC since $x(\theta^* + \varepsilon) = \bar{x}$. Likewise, if θ^* strictly preferred $\bar{x}$ to $\lim_{\theta \rightarrow \theta^* -} x(\theta)$, then by continuity, there is $\varepsilon > 0$ for which

$$U(x(\theta^* - \varepsilon), c(x(\theta^* - \varepsilon), \theta^* - \varepsilon), \theta^* - \varepsilon) < U(\bar{x}, E[c(\bar{x}, \theta) | \theta > \theta^*], \theta^* - \varepsilon)$$

also contradicting IC. $\square$

Lemma 13. *The graph of the solution ϕ to the IVP for price functions intersects L once at most.*

Proof. Suppose x^* is such that $(x^*, \phi(x^*)) \in L$. By construction, there is a type θ^* for which $\xi(\theta^*) = x^*$ and $c(\xi(\theta^*), \theta^*) = \phi(x^*)$. By construction of ϕ , $\phi(x) \geq q(x; \theta^*)$ for all x with equality holding when $x = x^*$. If there was another x^{**} such that $(x^{**}, \phi(x^{**})) \in L$ —and without loss of generality we may assume $x^{**} > x^*$ —then

since q is strictly decreasing in θ by definition, we would have $\phi(x) \geq q(x; \theta^*) > q(x; \theta^{**})$ and $\phi(x^{**}) = q(x^{**}; \theta^{**})$, a contradiction. $\square$

Proof of proposition 5. By proposition 3, there is no separating SAGE. Assume for contradiction that there is a pooling SAGE. Let $q(x; 0)$ be the indifference curve of type 0 passing through $(\bar{x}, E[c(\bar{x}, \theta)])$. By IC, $p(0) \geq q(0; 0)$, but by AG, Proposition 1.2, $p(0) \leq c(0, 0) < q(0; 0)$, via the assumption $U(\bar{x}, E[c(\bar{x}, \theta)], 0) < U(0, c(0, 0), 0)$. This yields the contradiction.

It follows that there is a semi-separating SAGE, containing a separating interval and a pooling interval. Again, by AG, Proposition 1.4, p is differentiable a.e., and wherever it is differentiable on the separating interval, it must satisfy $p' = -U_x/U_p$ lest the indifference curve of $\theta(x, p(x))$ crosses the graph of p , contradicting IC. By lemmas 8 and 9, it follows that p uniquely satisfies the IVP for price functions on the separating interval. To see that the separating interval is also (almost) unique, observe that if θ^* divides the separating and pooling intervals, then by lemma 12,

$$L \ni \lim_{\theta \rightarrow \theta^*} (x(\theta), c(x(\theta), \theta)) = \lim_{\theta \rightarrow \theta^*} (x(\theta), p(x(\theta))) \equiv \lim_{\theta \rightarrow \theta^*} (x(\theta), \phi(x(\theta)))$$

By lemma 13 that $\lim_{\theta \rightarrow \theta^*} (x(\theta), \phi(x(\theta))) = (\xi(\theta^*), c(\xi(\theta^*), \theta^*))$ for a unique θ^* . Thus the separating interval is either $[0, \theta^*)$ or $[0, \theta^*]$.

Lastly, the price function $p(x)$ for $x \in [\theta^*, \bar{\theta}]$ must trace along the indifference curve $q(x; \theta^*)$: any less and some type arbitrarily close to θ^* would strictly prefer x , violating IC, any more and no type is indifferent between $x(\theta)$ and x , violating AG, Proposition 1.2. $\square$

Lemma 14. *For a bound $b < \infty$ and a fixed ϕ_0 , the family*

$$\mathcal{F} = \{\phi \in AC[0, b] : \phi(x) = \phi_0 + \int_0^x f(s, \phi(s)) ds\}$$

is equicontinuous, and the functions $\underline{\phi}(x) := \inf_{\phi \in \mathcal{F}} \phi(x)$ and $\bar{\phi}(x) := \sup_{\phi \in \mathcal{F}} \phi(x)$

are in $\mathcal{F}$.

Proof. I will do the proof for $\underline{\phi}$, but the proof is the same for $\bar{\phi}$. By construction, $|f| \leq M$ for some $M < \infty$. So, for any $\varepsilon > 0$, we take $\delta = \varepsilon/2M$, then $\phi \in \mathcal{F}$ and $|x - y| < \delta$ implies

$$|\phi(x) - \phi(y)| = \left| \int_y^x f(s, \phi(s)) ds \right| \leq 2M|x - y| < \varepsilon$$

Thus $\mathcal{F}$ is equicontinuous. The infimum of an equicontinuous family is continuous. If $\{\phi_n\} \subset \mathcal{F}$ and $\phi_n \rightarrow \hat{\phi}$, then $\hat{\phi} \in \mathcal{F}$, because the Dominated Convergence Theorem implies

$$\hat{\phi}(x) = \lim_{n \rightarrow \infty} \phi_n(x) = \phi_0 + \lim_{n \rightarrow \infty} \int_0^x f(s, \phi_n(s)) ds = \phi_0 + \int_0^x f(s, \hat{\phi}(s)) ds$$

and the Fundamental Theorem of Calculus for Lebesgue integrals (Folland 1999) implies $\hat{\phi}$ is also absolutely continuous; thus proving $\hat{\phi} \in \mathcal{F}$. Moreover, it should be pointed out that $\hat{\phi}$ is continuously differentiable since f is continuous. Lastly, the Arzela-Ascoli Theorem implies $\mathcal{F}$ is sequentially compact.

At each x , we can define a sequence $\phi_n^x \in \mathcal{F}$ such that $\phi_n^x(x)$ decreases to $\underline{\phi}(x)$; there exists a subsequence which uniformly converges to a function ϕ^x , for which $\phi^x(x) = \underline{\phi}(x)$. One trivial consequence is that $\underline{\phi}$ is strictly monotonic and therefore a.e. differentiable, otherwise the graph of $\underline{\phi}$ and ϕ^x would cross, contradicting the definition of $\underline{\phi}$. Because $\underline{\phi}$ is both continuous and strictly monotonic on $[0, b]$, it must be absolutely continuous: otherwise, $\underline{\phi}' = -\infty$ on a positive measure set (Leoni 2017), and at x where $\underline{\phi}'(x) = -\infty$, the graph of ϕ^x would cross $\underline{\phi}$, by lemma 11, which contradicts the definition of $\underline{\phi}$. By the same logic, it must be the case that where $\underline{\phi}$ is differentiable, $\underline{\phi}'(x) = \phi_x^x(x) = f(x, \phi^x(x)) = f(x, \underline{\phi}(x))$. Since $\underline{\phi}$ is absolutely continuous and a.e. differentiable, it follows that $\underline{\phi} \in \mathcal{F}$. $\square$

Proof of proposition 6. f is continuous in both terms and is bounded on a bounded set, thus it satisfies the conditions of Carathéodory's existence theorem.

Consider once again the family

$$\mathcal{F} = \{\phi \in AC[0, b] : \phi(x) = \phi_0 + \int_0^x f(s, \phi(s)) ds\}$$

and the function $\bar{\phi}(x) := \sup_{\phi \in \mathcal{F}} \phi(x)$. Suppose $\mathcal{F}$ contains multiple functions, of which $\bar{\phi}$ is necessarily not Pareto dominant. Assume that the rest of the slope field north of $\text{graph}(\bar{\phi})$ is locally Lipschitz.

Recalling that $X = [0, \bar{x}]$, consider the SAGE associated with an environment with the restricted contract space $\tilde{X} = [x_*, \bar{x}]$. By the preceding propositions, a SAGE exists, and clearly, $\tilde{\phi} > \bar{\phi}$. Next, consider a sequence of such restricted environments $\tilde{X}_1 \subset \tilde{X}_2 \subset \dots$ such that $\bigcup_n \tilde{X}_n = [0, \bar{x}]$, each associated with a SAGE price function $\tilde{\phi}_n$. Since each $\tilde{\phi}_n$ is associated with a sequence of perturbations and a sequence of weak equilibria which converge to the SAGE, it follows there exists a sequence which converges to the SAGE associated with $\bar{\phi}$.

That EF selects for $\underline{\phi}$ follows from the uniqueness result in EF. □

Chapter 3

Fixed Costs, the Extensive Margin, and No-Trade

3.1 Introduction

It is curious that some insurance markets do not exist, and others regularly reject entire sub-populations of applicants. After all, even very risky groups of people are still risk-averse relative to an insurance company, so there should be gains to trade. Hendren (2013) comes to one surprising result: theoretically, no trade occurs only if the riskiest type $\bar{\theta}$ in the sub-population incurs a loss with probability $\theta = 1$, and empirically, a heavy-tailed distribution of losses appears to be a reasonable explanation in a number of markets, given non-trivial frictions. AG arrive at an even stronger theoretical result: in an endogenous contracts setting, no-trade occurs if and only if $\bar{\theta} = 1$.

But in some cases, that some people are highly certain of loss seems too stringent a condition for no-trade; here, adverse selection fails to fully explain rejections.¹² For example, why don't standard travel insurance policies cover those traveling to perform "high-risk" activities? Certainly, no one would travel abroad just to hurt

12. nor would moral hazard, a point first highlighted by Shavell (1979)

themselves, so buyers could hardly be *certain* of injury. One plausible explanation is that the fixed costs, such as underwriting and litigation, may be too high for general insurers to cover specialized risks. As Hendren (2017) alludes to in a model with moral hazard,¹³ adding fixed costs may render the no-trade condition sufficient *but not necessary*.

This chapter formalizes this idea in a competitive market with endogenous contracts whose key friction is a fixed cost κ . In doing so, I derive a necessary and sufficient condition for no-trade at D1 equilibrium (Banks and Sobel 1987) and describe the state of the market as it unravels towards no-trade. For clarity, the chronology of key results is as follows:

1. Without an *extensive* margin (choice of whether to exit the market), the equilibrium is fully separating along the *intensive* margin (choice of coverage *within* the market) when $\bar{\theta} < 1$.
2. For any given type defined by its probability of loss $\theta > 0$, the higher $\bar{\theta}$ is, the less coverage θ buys. When $\bar{\theta} = 1$, every other type buys 0 coverage. This is implied in Hendren's model.
3. Adding an extensive margin, with $\kappa > 0$ and $\bar{\theta} < 1$, leads low types to pool at 0 coverage.
4. The higher κ or $\bar{\theta}$ is, the more types pool at 0.
5. When $\bar{\theta}$ prefers buying 0 insurance to buying full insurance at cost, the market will have completely unraveled to no-trade.

To see that Hendren's necessary condition is a special case, observe that for someone who is certain of loss ($\theta = 1$), the variable cost of fully insuring against the loss would be the value of the loss itself, thus they weakly prefer 0 coverage and would

13. As we will see, we do not need or consider moral hazard, and in this particular example, one could even reason that there is sometimes little moral hazard after conditioning on high-risk activity.

strictly prefer it if insurance had a fixed cost $\kappa > 0$.

This chapter dovetails with the recent work by Geruso et al. (2023) characterizing a price-theoretic model of an insurance market with an extensive margin. Both this chapter and their paper assume perfect competition and appeal to the Riley (1979) concept. Theirs uses a framework with two fixed contracts (Akerlof 1970), whereas I use an endogenous contracts framework (Rothschild and Stiglitz 1976). They focus on analyzing policy tradeoffs, and while this model—with its fewer moving parts—may be well-suited to do so, I focus on illustrating the impact of frictions.

3.2 The Riley Model

3.2.1 Assumptions

The market consists of risk-neutral uninformed insurers¹⁴ and risk-averse buyers who seek to insure a loss $l > 0$ and who are informed of their probability θ of incurring the loss (their “type”), which Nature draws from a distribution F with support $\Theta = [0, \bar{\theta}]$ where $\sup \Theta = \bar{\theta} < 1$. For simplicity, we normalize the loss to 1 relative to buyers’ wealth to w . The utility to type θ of an insurance contract with coverage $x \in X = [0, 1]$ at price $p \leq x$, is given by

$$U(x, p, \theta) = \theta \cdot u(w - p - (1 - x)) + (1 - \theta) \cdot u(w - p)$$

where $u : \mathbb{R}_+ \rightarrow \mathbb{R} \in C^2$, $0 < u' < \infty$, $0 > u'' > -\infty$. The cost of supplying contract x to type θ is given by

$$c(x, \theta) = x\theta + \kappa$$

where $\kappa \geq 0$ is a fixed cost. For simplicity, I assume $w > 1 + \kappa$ so as to define u only over the positive reals. The fixed cost is immaterial for now and one may picture it

14. Specifically, a particular sub-population thereof as identified by observables like age, gender, etc.

to be 0.

The assumptions above are essentially the same as those in Hendren, bar two. First, we assume Θ is supported on an interval rather than an arbitrary compact set; this is simply convenient for ODE purposes, but not necessary, and moreover, from the empirical results of Hendren, it does not appear to be an errant assumption. Second, implicit in this model is that, for now, there is no individual rationality constraint; that is, no option to buy no insurance for no cost (except when $\kappa = 0$). We will soon relax this assumption.

Definition 4. The indifference curve of θ through (y, p) is a function $q^{y,p}(\cdot; \theta) : X \rightarrow \mathbb{R}$ defined implicitly by

$$U(x, q^{y,p}(x; \theta), \theta) - U(y, p, \theta) = 0,$$

which implies $q_x^{y,p}(x; \theta) = -U_x(x, q^{y,p}(x; \theta), \theta)/U_p(x, q^{y,p}(x; \theta), \theta)$.

Lemma 15. *Indifference curves are single-crossing on $x \in (0, 1], p \in [0, x]$.*

Proof. $-\frac{U_x(x, p, \theta)}{U_p(x, p, \theta)} = \frac{1}{1 + \frac{1-\theta}{\theta} \cdot \frac{u'(w-p)}{u'(w-p-(1-x))}}$ which is strictly increasing in θ . Thus if $\theta' > \theta$, then $q_x^{y,p}(y; \theta) < q_x^{y,p}(y; \theta')$, and so

$$\begin{aligned} q^{y,p}(x; \theta) - q^{y,p}(x; \theta') &= (q^{y,p}(x; \theta) - p) - (q^{y,p}(x; \theta') - p) \\ &= \int_y^x -\frac{U_x(x, q^{y,p}(x; \theta), \theta)}{U_p(x, q^{y,p}(x; \theta), \theta)} - \left(-\frac{U_x(x, q^{y,p}(x; \theta'), \theta')}{U_p(x, q^{y,p}(x; \theta'), \theta')} \right) dx \end{aligned}$$

< 0 if $x > y$ and > 0 if $x < y$. □

Lemma 16. $-\frac{U_x(x, p, \theta)}{U_p(x, p, \theta)} > c_x(x, \theta)$

Proof. We have

$$\frac{1}{1 + \frac{1-\theta}{\theta} \cdot \frac{u'(w-p)}{u'(w-p-(1-x))}} = \frac{\theta}{\theta + (1-\theta) \cdot \frac{u'(w-p)}{u'(w-p-(1-x))}} > \theta = c_x(x(\theta), \theta)$$

since $\frac{u'(w-p)}{u'(w-p-(1-x))} < 1$. □

3.2.2 Equilibrium

The Pareto-dominant, zero-profit separating equilibrium consists of an allocation x and a price p which solve the initial value problems¹⁵

$$x' = x \cdot \frac{\theta u'(w - 1 - \kappa + (1 - \theta)x) + (1 - \theta)u'(w - \theta x)}{\theta(1 - \theta)[u'(w - 1 - \kappa + (1 - \theta)x) - u'(w - \theta x)]} =: g(\theta, x) \quad x(\bar{\theta}) = 1 \quad (3.1)$$

$$p' = \frac{(p - \kappa)/x}{(p - \kappa)/x + (1 - (p - \kappa)/x) \cdot \left(\frac{u'(w-p)}{u'(w-p-(1-x))} \right)} =: h(x, p) \quad p(1) = c(1, \bar{\theta}) = \bar{\theta} + \kappa \quad (3.2)$$

These ODEs are derived from (Mailath and Von Thadden 2013)

$$\begin{aligned} x'(\theta) &= -\frac{U_p(x(\theta), c(x(\theta), \theta), \theta)}{U_x(x(\theta), c(x(\theta), \theta), \theta)} \cdot \frac{dc(x(\theta), \theta)}{d\theta} = -\frac{c_\theta(x(\theta), \theta)}{\frac{U_x(x(\theta), c(x(\theta), \theta), \theta)}{U_p(x(\theta), c(x(\theta), \theta), \theta)} + c_x(x(\theta), \theta)} \\ p'(x) &= -\frac{U_x(x, p(x), (p(x) - \kappa)/x)}{U_p(x, p(x), (p(x) - \kappa)/x)} \end{aligned}$$

where $\theta = \frac{p(x) - \kappa}{x}$ iff $c(x, \theta) = p(x)$; these in turn can be derived from the zero-profit condition and the first-order condition necessary for incentive compatibility:

$$\begin{aligned} p(x(\theta)) &= c(x(\theta), \theta) \\ \frac{dU(x(t), p(x(t)), \theta)}{dt} \Big|_{t=\theta} &= U_x(x(\theta), p(x(\theta)), \theta) \cdot x'(\theta) \\ &\quad + U_p(x(\theta), p(x(\theta)), \theta) \cdot p'(x(\theta)) \cdot x'(\theta) = 0. \end{aligned}$$

The Pareto-dominant, zero-profit separating equilibrium is also a *reactive* equilibrium (Riley 1979; Engers and Fernandez 1987); loosely speaking, for any profitable attempt at cream skimming the market, there exists a market *reaction* which renders the attempt unprofitable. We refer to this equilibrium as the Riley equilibrium,

15. I will use x (the variable) to denote an arbitrary contract with coverage $x \in X$, and $x : \Theta \rightarrow X$ (the function) to denote an allocation. The context will clarify which is which.

which the following statements concern.

Proposition 17. $x'(\theta) =: g(\theta, x(\theta)), x(\bar{\theta}) = 1$ has a unique solution on $D_x = \{(\theta, x) : \theta \in (0, \bar{\theta} + \varepsilon), x \in (-\varepsilon, 1 + \varepsilon)\}$, for a sufficiently small ε . Likewise, $p'(x) =: h(x, p(x)), p(1) = \bar{\theta}$ has a unique solution on $D_p = \{(x, p) : x \in (0, 1 + \varepsilon), p \in (\kappa - \varepsilon, \kappa + x)\}$, for a sufficiently small ε .

Proposition 18. Setting $x(0) = 0$, x is a (continuous, strictly increasing) bijection from Θ to $[0, 1]$; for each θ , there is a unique x for which $p(x) = c(x, \theta)$,

Proposition 19. The equilibrium is incentive-compatible (IC).

Proposition 20. As the support of the distribution approaches $[0, 1]$ (that is, as $\bar{\theta} \rightarrow 1$), the Riley price function approaches $x \mapsto x + \kappa$ uniformly.

Corollary 20.1. Let $\bar{\theta}_n \uparrow 1$, and let (x_n, p_n) be the associated Riley equilibrium. For any fixed $\theta < 1$, $x_n(\theta) \rightarrow 0$.

Proposition 17 states that the equilibrium is unique. Proposition 18 points out the possibly surprising fact that, regardless of what $\bar{\theta}$ is, $\lim_{\theta \rightarrow 0} x(\theta) = 0$; *id est*, not only is there unraveling along the intensive margin, but types are ‘spread out’ over whole of $X = [0, 1]$. Corollary 20.1 is the Riley analog of the necessary condition in Hendren (2013), which in Azevedo and Gottlieb (2017) is sufficient.

There are many ways to understand this result, and I see it this way: the Riley price must be a) zero-profit and b) IC. a) implies $p(x(\theta)) = c(x(\theta), \theta)$, and b) implies $q^{x(\theta), p(x(\theta))}(x; \theta) \leq p(x)$. Unfortunately, the indifference curve of $\theta = 1$ is exactly its cost curve; thus if $\bar{\theta} = 1$, then $p(x) \geq c(x, 1) = x + \kappa$. A closely aligned strand of reasoning is that in Riley equilibrium—and in the D1 equilibrium considered in the next section—pools are unstable because of cream skimming, so the market must separate the types; but as $\bar{\theta}_n \uparrow 1$, it becomes harder and harder to do that, and when $\bar{\theta} = 1$, it becomes impossible to separate.

3.3 D1 Equilibrium

In effect, the above described model is a signaling game in which the buyer sends signal x , the seller with belief μ responds mechanically with $p(x) = E_{t \sim \mu(x, \cdot)}[c(x, t)]$, leaving the buyer with utility $U(x, p(x), \theta)$. That motivates the following definition analogous to a Perfect Bayesian Equilibrium.

Definition 5. (x, μ) is an *equilibrium* if:

1. μ is a regular conditional distribution (RCD) of θ given $x \circ \theta$; that is, μ satisfies Bayes' Rule
2. $x(\theta) \in \arg \max_{y \in X} U(y, E_{t \sim \mu(y, \cdot)}[c(y, t)], \theta)$ for all $\theta \in \Theta$; that is, x is IC

From this we simply define the distribution and conditional expectations as

$$F(\theta|x) = \mu(x, \cdot)$$

$$E[f(x, \theta)|x] = \int \mu(x, d\theta) f(x, \theta)$$

for any measurable function f .

Definition 6. (x, μ) satisfies the *D1 criterion* if for all $y \in X \setminus x(\Theta)$, $\mu(y, \cdot)$ satisfies the following: if θ', θ'' satisfies

$$\begin{aligned} U(y, p, \theta') &\geq U(x(\theta'), E[c(x(\theta'), t)|x(\theta')], \theta') \\ \implies U(y, p, \theta'') &> U(x(\theta''), E[c(x(\theta''), t)|x(\theta'')], \theta'') \end{aligned} \tag{3.3}$$

for all p , then $\theta' \notin \text{supp } F(\theta|y)$.

In lay terms, this is to say that when formulating beliefs about who might have deviated to an off-the-equilibrium action y , θ' is eliminated if there is a type θ'' which satisfies the following: for any (y, p) which θ' weakly prefers to its equilibrium action, θ'' *strictly* prefers (y, p) to *its* equilibrium action.

The D1 equilibrium typically aligns with the Riley (1979) and AG equilibria. However, neither model permits a discontinuous cost function, which we turn to.

3.4 Adding the Extensive Margin

All assumptions in section 3.2.1 are maintained except we augment the cost to include a fixed cost $\kappa \geq 0$ whenever positive coverage is provided:

$$c(x, \theta) = x\theta + \kappa \mathbb{1}_{x>0}.$$

In this sense, the “null contract” $x = 0$ for which $c(x, \theta) = 0$ for all θ constitutes an extensive margin, that is, “not buying” is the same as “buying nothing for the price of nothing.” Thus, in this model, IC implies individual rationality.

The Riley allocation x_r is the reactive equilibrium in the event that there is no extensive margin (that is, choosing $x = 0$ at $p = 0$ is not an option), as defined in the preceding sections. We denote the associated Riley price function p_r and Riley allocation x_r . Under the following assumption, the succeeding propositions show that the unique D1 equilibrium involves high risk types being allocated to x_r .

Assumption 5. $U(1, \bar{\theta} + \kappa, \bar{\theta}) > U(0, 0, \bar{\theta})$; the riskiest type prefers paying an actuarially fair price for full insurance to no insurance at all.

Proposition 21. *Given assumption 5, there is a unique θ for which*

$$U(x_r(\theta), p_r(x_r(\theta)), \theta) = U(0, 0, \theta).$$

Denoting this type as $\hat{\theta}$, all types $> \hat{\theta}$ strictly prefer the Riley allocation over the null allocation given actuarially fair prices, and the reverse is true for types $< \hat{\theta}$.

Proposition 22. *Let x be an allocation satisfying*

$$x(\theta) = \begin{cases} 0 & (0 \leq \theta < \hat{\theta}) \\ 0 \text{ or } x_r(\theta) & (\theta = \hat{\theta}) \\ x_r(\theta) & (\hat{\theta} < \theta \leq \bar{\theta}) \end{cases}$$

and let $\mu : X \times \mathcal{B}_\Theta \rightarrow [0, 1]$ be a regular conditional distribution of θ given $x \circ \theta$

where, for each off-equilibrium $y \in X \setminus x(\Theta)$, $\mu(y, A) = \mathbb{1}_{\bar{\theta} \in A}$. Given assumption 5, (x, μ) is a D1 equilibrium.

Definition 7. y is a pool of (x, μ) if $y \in x(\Theta)$ and $\mu(y, \cdot)$ is not a degenerate distribution.

Proposition 23. If $U(1, \bar{\theta} + \kappa, \bar{\theta}) > U(0, 0, \bar{\theta})$, i.e. assumption 5 is satisfied, there are no D1 equilibria except the one described in proposition 22. If $U(1, \bar{\theta} + \kappa, \bar{\theta}) < U(0, 0, \bar{\theta})$ (resp. $\leq$), then all types (resp. all but possibly $\bar{\theta}$) pool at 0 in D1 equilibrium.

The results and their proofs are standard for D1 equilibria under single crossing preferences. Namely, any off-the-equilibrium path contract purchased with coverage less than a nearby pool should be assigned almost surely to the infimum of the types in the pool, and any contract greater than the nearby pool should be assigned to the supremum.

3.5 Concluding Remarks

Because of the distribution-free nature of least-cost separating equilibria, a type space (where types are defined by their probability of loss) with a large upper bound $\bar{\theta}$ can result in significant Rothschild-Stiglitz-style unraveling, such that small frictions can lead to the non-existence of the market. For a model where said friction is a fixed cost, this chapter has arrived at a simple condition for no-trade in D1 equilibrium: $U(1, \bar{\theta} + \kappa, \bar{\theta}) < U(0, 0, \bar{\theta})$.

3.6 Chapter Appendix: Omitted Proofs

Proof of proposition 17. u' is C^1 on an open set containing $[w - 1 - \kappa, w]$, so $g(\theta, \cdot)$ is composed of C^1 functions on $(-\varepsilon, 1 + \varepsilon)$ for sufficiently small ε for any fixed $\theta \in (0, \bar{\theta} + \varepsilon)$, which implies g is locally Lipschitz in x .¹⁶ Likewise, $h(x, \cdot)$ is composed

16. In particular, for fixed θ , the u' terms in the numerator are C^1 so denote the numerator $a(x) \in C^1$. The ones in the denominator are also C^1 , and the denominator is bounded below by a

of C^1 functions on $(-\varepsilon, x)$ for any fixed $x \in (0, 1 + \varepsilon)$. The result follows from the Picard–Lindelöf theorem (Hale 1980). $\square$

Proof of proposition 18. By construction, x is strictly increasing on its interval of existence. Plugging in $y(\theta) = 0$ shows that $y'(\theta) = g(\theta, y(\theta))$. By uniqueness on D_x , $x(\theta) > 0$ for all $\theta > 0$; that is, there does not exist $\theta > 0$ for which $x(\theta) = 0$, this shows its maximal interval of existence on D_x includes Θ .

By inspection of g , it can be seen that x is thereby strictly monotonic for $\theta \in (0, \bar{\theta}]$. Similar to before, $q(x) = \kappa$ solves $q_x(x) = h(x, q(x))$, therefore, $p(x) > \kappa$ whenever $x \in (0, 1]$. Assume for contradiction that $\lim_{\theta \rightarrow 0} x(\theta) = c \in (0, 1]$. This implies $\lim_{\theta \rightarrow 0} p(x(\theta)) = p(c) > \kappa$. However, $\lim_{\theta \rightarrow 0} c(x(\theta), \theta) = \lim_{\theta \rightarrow 0} \theta x(\theta) + \kappa = \kappa$, and a contradiction is drawn. It follows $\lim_{\theta \rightarrow 0} x(\theta) = 0$, and it is continuous to define $x(0)$ as such. $\square$

Proof of proposition 19. The utility difference to θ between $x(\theta)$ and another contract $x(\hat{\theta})$ is

$$\begin{aligned} & U(x(\theta), p(x(\theta)), \theta) - U(x(\hat{\theta}), p(x(\hat{\theta})), \theta) \\ &= \int_{\hat{\theta}}^{\theta} \frac{dU(x(t), p(x(t)), \theta)}{dt} dt \\ &= \int_{\hat{\theta}}^{\theta} U_x(x(t), p(x(t)), \theta) \cdot x'(t) + U_p(x(t), p(x(t)), \theta) \cdot p'(x(t)) \cdot x'(t) dt \\ &= \int_{\hat{\theta}}^{\theta} x'(t) [U_x(x(t), p(x(t)), \theta) + U_p(x(t), p(x(t)), \theta) \cdot p'(x(t))] dt \\ &= \int_{\hat{\theta}}^{\theta} x'(t) \left[U_x(x(t), p(x(t)), \theta) - U_p(x(t), p(x(t)), \theta) \cdot \frac{U_x(x(t), p(x(t)), t)}{U_p(x(t), p(x(t)), t)} \right] dt. \end{aligned}$$

x' is always positive. The single-crossing condition implies the rest of the integrand is positive iff $\theta > \hat{\theta}$ and negative iff $\theta < \hat{\theta}$, which ensures the definite integral is positive. Ergo, $x(\theta)$ is an optimal strategy for θ . $\square$

number > 0 , so denote the denominator as $b(x) \in C^1$ with $\inf b > 0$. $a(x)/b(x)$ is thus C^1 .

Proof of proposition 20. Clearly, on $[0, 1]$, $p(x) = x + \kappa$ uniquely solves the IVP

$$p'(x) = h(x, p(x)) \quad p(1) = 1 + \kappa$$

Let $\{\theta_n\} \subset (0, 1)$ be a sequence increasing to 1. Let p_n solve

$$p'_n(x) = h(x, p_n(x)) \quad p_n(1) = c(1, \bar{\theta}_n) = \bar{\theta}_n + \kappa$$

$\theta_n \rightarrow 1$ implies $c(1, \bar{\theta}_n) \rightarrow 1 + \kappa$. Thus, p_n converges uniformly to the function which solves

$$\phi'(x) = h(x, \phi(x)) \quad \phi(1) = 1 + \kappa$$

(Teschl 2012, Theorem 2.8). Ergo, $p_n(x) \rightarrow p(x) = x + \kappa$ uniformly. $\square$

Proof of corollary 20.1. Fix $\theta < 1$. Assume for contradiction that there is $\varepsilon > 0$ such that $x_n(\theta) > \varepsilon$ for infinitely many n . By assumption, $p_n(x_n(\theta)) = c(x_n(\theta), \theta) = \theta x_n(\theta) + \kappa$. Consequently, there are infinitely many n for which $x_n(\theta) + \kappa - p_n(x_n(\theta)) = x_n(\theta) + \kappa - \theta x_n(\theta) - \kappa = (1 - \theta)x_n(\theta) > (1 - \theta)\varepsilon$. Ergo, $p_n(x)$ does *not* converge uniformly to $x + \kappa$. $\square$

Proof of proposition 21. The difference in utility between the Riley allocation and the null allocation for type θ is

$$\delta(\theta) = U(x_r(\theta), \theta x_r(\theta) + \kappa, \theta) - U(0, 0, \theta)$$

By assumption 5, $\delta(\bar{\theta}) > 0$, and since $x_r(0) = 0$, we also see that $\delta(0) < 0$; so since δ is continuous, by the Intermediate Value Theorem, there exists at least one $\hat{\theta} \in (0, \bar{\theta})$ where $\delta(\hat{\theta}) = 0$.

I now show that $\hat{\theta}$ is unique by showing that $\delta(t) < 0$ if $t < \hat{\theta}$ and $\delta(t) > 0$ if

$t > \hat{\theta}$. Let $\hat{x} = x(\hat{\theta})$ and $\hat{p} = p(x(\hat{\theta}))$. The indifference curve $q^{\hat{x}, \hat{p}}(x; \theta)$ solves

$$q^{\hat{x}, \hat{p}}(x; \theta) = \hat{p} + \int_{\hat{x}}^x -\frac{U_x(y, q^{\hat{x}, \hat{p}}(y; \theta), \theta)}{U_p(y, q^{\hat{x}, \hat{p}}(y; \theta), \theta)} dy$$

Let $t > \hat{\theta}$. Single-crossing (see lemma 15) implies $q^{\hat{x}, \hat{p}}(x; t) < q^{\hat{x}, \hat{p}}(x; \hat{\theta})$ whenever $x < \hat{x}$, so there exists $x^* > 0$ for which $q^{\hat{x}, \hat{p}}(x^*; t) = q^{\hat{x}, \hat{p}}(0; \hat{\theta}) = 0$, thus we have

$$U(0, 0, t) < U(x^*, 0, t) = U(\hat{x}, \hat{p}, t) \leq U(x_r(t), p_r(x_r(t)), t)$$

where the latter inequality follows by IC, as needed. Now, let $t < \hat{\theta}$. By IC, $q^{\hat{x}, \hat{p}}(x; \hat{\theta}) \leq p_r(x)$. By single-crossing, $q^{0,0}(x; t) < q^{0,0}(x; \hat{\theta}) = q^{\hat{x}, \hat{p}}(x; \hat{\theta}) \leq p_r(x)$ for $x > 0$, and therefore,

$$U(x_r(t), p_r(x_r(t)), t) > U(x(t), q^{0,0}(x(t); t), t) = U(0, 0, t)$$

as needed. □

Proof of proposition 22. (x, μ) is IC by proposition 21. By construction, for $y \in X \setminus x(\Theta)$, $\text{supp } F(\theta|y) = \{\hat{\theta}\}$. Lastly, we need to check that the implication eq. (3.3) is never satisfied for $\hat{\theta}$. Let θ'' be some other type and let $y \in X \setminus x(\Theta)$. Let q once again denote the indifference curve. By definition

$$U(y, q^{x(\hat{\theta}), c(x(\hat{\theta}), \hat{\theta})}(y; \hat{\theta}), \hat{\theta}) \geq U(x(\hat{\theta}), c(x(\hat{\theta}), \hat{\theta}), \hat{\theta})$$

By single-crossing, if $\theta'' < \hat{\theta}$, then $q^{0,0}(y, \theta'') < q^{0,0}(y, \hat{\theta})$, so

$$U(0, 0, \theta'') > U(y, q^{x(\hat{\theta}), c(x(\hat{\theta}), \hat{\theta})}(y; \hat{\theta}), \theta'')$$

Conversely, if $\theta'' > \hat{\theta}$, θ'' then, by IC and single-crossing,

$$\begin{aligned} U(x_r(\theta''), c(x_r(\theta''), \theta''), \theta'') &\geq U(x_r(\hat{\theta}), c(x_r(\hat{\theta}), \hat{\theta}), \theta'') \\ &= U(y, q^{x_r(\hat{\theta}), c(x_r(\hat{\theta}), \hat{\theta})}(y, \theta''), \theta'') \\ &> U(y, q^{x_r(\hat{\theta}), c(x_r(\hat{\theta}), \hat{\theta})}(y, \hat{\theta}), \theta'') \end{aligned}$$

as needed. $\square$

Lemma 24. *If (x, μ) is IC and $y \in X$ is a pool then $x^{-1}(y)$ is an interval of Θ .*

Proof. Let y be a pool and let $\theta < \theta' < \theta''$. By single-crossing,

$$\max\{q^{y, p(y)}(x; \theta), q^{y, p(y)}(x; \theta'')\} > q^{y, p(y)}(x; \theta')$$

for all $x \neq y$. So anything θ' weakly prefers to y is strongly preferred by either θ or θ'' . $\square$

Lemma 25. *If (x, μ) is IC then x is increasing. It is constant on pooling intervals (obviously) and strictly increasing on separating intervals.*

Proof. Suppose for contradiction that x was not strictly decreasing on a separating interval T . Since x is not pooling on T , this implies there is θ', θ'' for which $\theta' > \theta''$ but $x(\theta') < x(\theta'')$. For $x > x(\theta')$, $q^{x(\theta'), c(x(\theta'), \theta')}(x; \theta') > c(x, \theta')$ by lemma 16; therefore,

$$q^{x(\theta'), c(x(\theta'), \theta')}(x(\theta''); \theta') > c(x(\theta''), \theta') > c(x(\theta''), \theta'')$$

which implies $U(x(\theta'), c(x(\theta'), \theta'), \theta') < U(x(\theta''), c(x(\theta''), \theta''), \theta')$; that is, θ' prefers $x(\theta'')$, which contradicts IC. $\square$

Lemma 26. *If (x, μ) is a D1 equilibrium then there is no pool $y > 0$.*

Proof. Suppose for contradiction that there was a pool at $y > 0$, and let $\theta_* = \inf x^{-1}(y)$, $\theta^* = \sup x^{-1}(y)$. Several facts follow. First, by the single-crossing condition, if θ_* does not choose y , then it must be indifferent between $x(\theta_*)$ and y , lest some

type arbitrarily close would also strictly prefer $x(\theta^*)$. Second, because $\mu(y, \cdot)$ is non-degenerate, there must be an interval $(\theta_*, \hat{\theta}) \subset x^{-1}(y)$ for whom $c(y, \theta)$ is less than $p(y) = E[c(y, t) | t \in x^{-1}(y)]$. Third, the interval $I = \{x < y : q^{y, p(y)}(x; \theta_*) > c(x, \theta_*)\}$ where the indifference curve of θ_* passes over its cost curve must contain contracts which are out-of-equilibrium. Assume for contradiction $x(\theta) \in I$. If $\theta > \theta^*$, then $p(x) \geq c(x, \theta^*)$, in which case by single-crossing, θ would strictly prefer y to x . If $\theta < \theta_*$, then $p(x) \leq c(x, \theta_*)$, and by definition of I , θ_* would strictly prefer x to y .

In D1, for all $x \in I$, $\mu(x, A) = \mathbb{1}_{\theta_* \in A}$ (the Dirac measure for θ_*). But then $x(\theta_*)$ is not IC, because $c(\cdot, \theta_*) < q^{y, p(y)}(x; \theta_*)$ on I , so θ_* would be better off going off-the-equilibrium. $\square$

Proof of proposition 23. x is not separating since some type θ arbitrarily close to 0 would strictly prefer $(0, 0)$ to $(x(\theta), c(x(\theta), \theta))$. Any semi-separating/pooling equilibrium pools only at 0. Mailath and Von Thadden (2013, Theorem 3) implies that on the separating portion of any semi-separating equilibrium, $x = x_r$. By proposition 21, the semi-separating interval is either $(\hat{\theta}, 1]$ or $[\hat{\theta}, 1]$, for a unique $\hat{\theta}$. This concludes the proof that x must be as in proposition 22, given assumption 5. As for beliefs, the D1 criterion and single-crossing jointly imply $\mu(y, A) = \mathbb{1}_{\hat{\theta} \in A}$ for $y \in X \setminus x(\Theta)$. Trivially, if $U(1, \bar{\theta} + \kappa, \bar{\theta}) = U(0, 0, \bar{\theta})$, then $\hat{\theta} = \bar{\theta}$, and if $U(1, \bar{\theta} + \kappa, \bar{\theta}) < U(0, 0, \bar{\theta})$, then $q_{0,0}(x; \bar{\theta}) < p_r(x)$ —no part of the Riley allocation is IC and a pool forms at 0. $\square$

Chapter 4

Dynamic Stochastic Selection Markets and the Implications of Inattention

4.1 Introduction

The empirical literature on consumer choice in selection markets has burgeoned in the last few decades,¹⁷ the implications of which motivate, and provide the basis for, new theoretical developments. Case in point, classical selection market models tend to assume one dimension of private information, so when empirical evidence suggested the importance of *multiple* dimensions (Finkelstein and McGarry 2006), this motivated the development of selection models with multidimensional types (Veiga and Weyl 2016; Azevedo and Gottlieb 2017). Though these models do not supplant classical intuition (Friedman 1962; Akerlof 1970; Rothschild and Stiglitz 1976), they do complement it by allowing us to understand where different assumptions lead, and to square observations in the real world with the model world.

17. Recent handbook chapters include: Chandra, Handel, and Schwartzstein (2019), Einav, Finkelstein, and Mahoney (2021), and Handel and Ho (2021)

To wit, this chapter aims to provide a tractable approach to incorporating choice frictions in a dynamic setting, and bridge the gap between existing models and three important stylized facts about selection markets:

1. Agents change decisions slowly over time (e.g. Handel (2013)), while classical models focus on static equilibria
2. Agents make inconsistent choices (e.g. Abaluck and Gruber (2011)), while classical models assume consumers optimize in equilibrium
3. Some markets exist and some markets don't (e.g. Hendren (2017)), but markets display little of their predicted fragility¹⁸

Beyond explanation, a secondary motivation which informs the construction of the model presented herein is to make substantive predictions. Empirical work in selection markets demonstrate that statements about them are often context dependent: better information *could* lead to better health plan choices, but it could also lead to greater unraveling, so the welfare effect could go either way. A tractable computational model could be a useful tool for informing policy in such a setting.

This chapter provides two contributions: first, I set out a selection market framework which allows for i) a time dimension, ii) micro-founded stochastic choices, and iii) multiple dimensions of information and contracts, including moral hazard; a dynamic informational stochastic equilibrium (DISE), for short. Second, I present the *linear logit model*—the simplest model within this framework capable of capturing choice frictions; I use the model to isolate the implications of rational inattention in active choices in selection markets.¹⁹

18. I am referring here to the persistent critique of models of both strategic interactions (Cho and Kreps 1987; Banks and Sobel 1987) and markets (Riley 1979; Azevedo and Gottlieb 2017) which imply least-cost separation where possible: that an arbitrarily small changes to the model parameters can cause significant changes in the predicted equilibrium.

19. *Passive* choice refers to whether to re-consider one's action, *active* choice refers to what action to take next.

Despite its simplicity, the benchmark model elicits some surprising implications when compared beside classical results. First, it reproduces the familiar Riley outcome found in signaling (Cho and Kreps 1987; Banks and Sobel 1987) and selection markets (Handel, Hendel, and Whinston 2015; Azevedo and Gottlieb 2017), only with added noise. Second, using a single attention parameter $\beta \in [0, \infty]$ which can be interpreted as the inverse of an information cost, we show that the unstable theoretical nature of Riley outcomes—whereby a small change in the type distribution can cause large market changes—is reliably reproduced only when β is large. That is, the noise added by β can collapse least-cost separations into pools. Intuitively, if good types vastly outnumber bad types, and some proportion of the good types fail to separate even when incentivized, that small group of good types may still swamp the bad types, thus reducing the incentive to separate; as the incentive to separate gets smaller and smaller, this collapses the separation towards a pool. Third, along similar lines, I show that welfare can be but is not necessarily improved by allowing agents to make better choices; intuitively, agents can gain by *either* improving their individual choices *or* minimizing costly separation, sometimes creating a tradeoff in β . Lastly, following Azevedo and Gottlieb (2017), I calibrate a health insurance model with 26 evenly-spaced contracts and 96,000 types to the parameter estimates of Einav et al. (2013).

4.2 Background and Related Literature

*Theory of adverse selection markets.*²⁰ Two seminal contributions form the starting point for analyzing selection markets: Akerlof (1970, hereafter, *Lemons*) and Rothschild and Stiglitz (1976, hereafter, RS). *Lemons* describes a market with a *fixed* contract space. With two contracts with coverage $\{L, H\} \subset [0, 1]$, the first-best world has everyone buy high coverage H , but in the adverse selection world, only the most risky—so-called “lemons”—buy H because they have the highest willingness to pay, thus inflating its price and deflating quantity traded to suboptimal levels.

20. A more comprehensive up-to-date overview of both theory and empirics can be found in Einav, Finkelstein, and Mahoney (2021).

H is said to *unravel*. The case with more contracts is more complicated, but the suboptimality holds. In the limit, with an interval of contracts, we get an RS model in which insurers can choose what contracts in $[0, 1]$ to offer. RS argues that in such a market, no Nash equilibrium exists. Indeed under some model specifications, any potential equilibrium could be disturbed by offering slightly different coverage at a price which attracts the right types.²¹

RS is typically presented as a *screening* game in which the uninformed players move first (by offering contracts), and informed players move second. A parallel literature on *signaling* games—in which informed agents move first—began with the seminal work of Spence (1973), who proposed that non-lemons could take costly actions which lemons find more costly in order to “separate.” Equilibria in Spencian games—consisting of *signals* from the informed and (Bayesian) *beliefs* of the uninformed—came later to be formalized as Perfect Bayesian Equilibria (PBE).²² In insurance, the costly action which separates low risk types would be buying a lower coverage contract. The problem with signaling is, however, that rather than too few equilibria (as in RS), there are infinitely many PBE.

While the distinction between signaling and screening is important in strategic settings (i.e. games), it is less so when considering markets where thousands of uninformed and informed agents move every period, with prices constantly reinforcing beliefs and vice versa. What’s more important is choosing an equilibrium concept (ideally, yielding > 0 and $< \infty$ equilibria). In endogenous contracts and in signaling, one equilibrium outcome has proven influential—the least-cost separating equilibrium, better known as the *Riley outcome*—which Riley (1979) showed coincides with the *reactive equilibrium*, consisting of a price function over signals where any attempt at cream skinning could be met by a profitable *reaction* by the market which renders the skim unprofitable. Later, equally influential work in *belief-based refinements* of

21. Hendren (2014) contrasted the two thusly: *Lemons* describes an “equilibrium of market unraveling” and RS describes an “unraveling of market equilibrium.”

22. Fudenberg and Tirole (1991) say the first use of PBE was in Milgrom and Roberts (1982)

games, which restricted the set of plausible PBE by eliminating certain beliefs, also pointed towards a Riley outcome (Cho and Kreps 1987; Banks and Sobel 1987); as has recent work introducing a price-taking equilibrium concept for selection markets which allows for multidimensional type and contract spaces (Azevedo and Gottlieb 2017, hereafter, AG). The Riley outcome, however, suffers from a lasting weakness: namely, that the outcome depends only on the support of the type distribution, and not the distribution itself. One consequence is that good types separate from bad types *regardless* of how few bad types there are, as long as it is not zero.²³ Some recent empirical work also use distribution-free results: Hendren (2013) and Hendren (2017) are based partly off the premise that a necessary condition for market collapse is that the support contains a very risky type.

This chapter differs from all the above mentioned work on two fronts. One, previous work focused on static equilibria (or lack thereof), while this work focuses on dynamics. Second, previous work assumed agents made utility-maximizing choices of contracts, whereas I assume stochastic choice micro-founded by rational inattention. Both prove crucial in the results to come.

Stickiness and inattention. This work contributes to a rich economic literature emphasizing stickiness and inattention as crucial choice frictions. However, the literature on which our models are based has concerned not selection markets, but price-setting behaviour in macroeconomics. After the sticky markets of the neo-classical synthesis (Keynes 1936; Hicks 1937) gave way to micro-founded models, stickiness re-emerged in the form of New Keynesian models which have come to

23. This has been pointed out by AG and by Mailath, Okuno-Fujiwara, and Postlewaite (1993). The latter offers an alternative equilibrium concept, the *undefeated equilibrium* (UE), in which high productivity workers do not separate from low productivity workers if the separating PBE is worse for them than the pooling PBE. This happens when the low productivity workers are sufficiently few in number that wages for high type workers would not be sufficiently depressed to warrant separating. This is an important alternative to Riley outcomes in signaling games, but in a market in which workers enter and take wage scales as given, there does not appear to be a mechanism by which a separating equilibrium would collapse to a pool (or vice versa) were the distribution of types to shift from one conducive to a pooling UE versus one conducive to a separating UE.

dominate macroeconomics. The seminal contribution of Calvo (1983) considered a simple model where firms had some probability each period to be able to reset prices. This proved surprisingly effective for as simple a model as it was; nonetheless, stories still differed on why prices were sticky. It was in this context that a research agenda focusing on rational inattention came to be (Sims 1998, 2003). An important contribution has been Matějka and McKay (2015), who connected the multinomial logit stochastic choice model—a tractable model commonly used in discrete choice econometrics—to rational inattention. For its part, this chapter studies the effects of similar choice frictions on selection markets, focusing on the most simple case in which agents move on from old decisions at an exogenous rate (as in the Calvo model) and make new choices via multinomial logit (as in the Matějka-McKay model). But similar to the conservative interpretations that much of the rational inattention literature takes on the topic, I would interpret the model herein to be an *as-if* model, rather than how agents really act.²⁴

Related work. Two recent empirical works have leveraged the Matějka-McKay model as a microfoundation for a logit choice model of selection markets: Brown and Jeon (2023) and the job market paper of Boehm (2024). I am not aware, however, of any similar work of a primarily theoretical nature.

4.3 Theoretical Framework

4.3.1 Basic Setup

The goal of this section is not to set out a functional model *per se*, but to outline the key desiderata of dynamic informational stochastic equilibria—DISE²⁵—without necessarily incorporating selection.

24. For example, in insurance, this chapter models a world where it is *as-if* a certain percentage of people move on from their old contracts each year and were rationally inattentive in picking new ones.

25. unfortunately, leaving out ‘stochastic’ makes the acronym needlessly morbid

Setting. The model is set in discrete time. There is a finite *contract space*²⁶ X , two finite type spaces $\Theta_1 \cup \Theta_2 = \Theta$ representing the private information held on either side of the market (termed “buyers” and “sellers”). To each type $\theta \in \Theta$ is a prior probability $\mu(X, \theta)$ defining the mass of each type over their respective spaces. The total mass of buyers $\mu(X, \Theta_1)$ is normalized to one.

Outflows and Inflows. To each type θ , contract x , and time $t \in \mathbb{N}$ is a corresponding mass $\mu^t(x, \theta)$. In a *fixed types* model, $\sum_x \mu^t(x, \theta) = \mu^t(X, \theta)$ is constant over all t for any fixed θ ; when there is no confusion, the time index is dropped. At each time, to each (x, θ) , there is an outflow $\nu^t(x, \theta)$ and an inflow $\iota^t(x, \theta)$ such that $\mu^{t+1} - \mu^t = \iota^t - \nu^t$. Of course, for any θ , total inflow and outflow must equate: $\iota(X, \theta) = \nu(X, \theta)$.

Pricing mechanism. There is a specified means by which prices in t are set based on the state variables and choices made at $t - 1$, *including for contracts which are not traded*.

Preferences. There exists a utility function $U(x, p, \theta)$ which is strictly monotonic in p .

Stochastic choices. $\iota > 0$ for all (x, θ) and $\nu > 0$ whenever $\mu > 0$. This is to say that θ has a positive, however small, probability of choosing/leaving any given contract x . Ideally, $\iota, \nu \ll 1$. The outflow and/or the inflow should be strictly monotonic in U ; that is, types should be more likely to flow away from contracts they do not prefer and/or towards contracts they do prefer.

Equilibrium. If it exists, an *inflow-outflow equilibrium* or “steady-state” is achieved when outflows equal inflows for almost every (x, θ) .

Model flexibility. The framework leaves a lot of room for modeling decisions. In particular, choice frictions can be incorporated into the outflows and inflows, with

26. or signal space, action space, etc.

the outflows typically capturing the “passive” choice frictions and inflows capturing the “active” choice frictions.

There is some flexibility in how the model sets its prices, but each model typically restricts the set of possible price setting mechanisms. A search-and-matching model, for example, typically specifies how surplus is split. Models with one-sided information asymmetry typically require a mechanism specifying how the uninformed side of the market might use the contract space (provided it is sufficiently large) to cream skim the informed side.

Intuitively, one of the things that makes this model work in practice is that *not everybody moves at once*. “*If everyone were to move at once, what should each person do?*” is a much more difficult question than “*If an infinitesimal mass of people were to move conditional on no one else moving, what are some good moves?*” because the movement of the infinitesimal mass has no effect on the state of the market. As we will see later, the stochastic nature of the flows are helpful in that equilibria tend to be more stable when people make mistakes that are random on a micro-level but predictable on a macro-level.

4.3.2 Linear Logit Model

In the benchmark model, only one side of the market moves and is privately informed, termed the “agents.” The other side consists of arbitrarily many homogeneous, uninformed, risk-neutral, competitive firms that, in competition, make zero profit. Agents come with a cost c or a value V , depending on their privately informed type and chosen contract.

The outflow is linear in mass

$$\nu(x, \theta) = \alpha \mu(x, \theta) \tag{4.1}$$

where α is a rate parameter, and the inflow

$$\iota(x, \theta) = \rho(x, \theta) \cdot \nu(X, \theta)$$

is given by the total outflow times a multinomial logit (better known in other fields as *softmax*) probability

$$\rho(x, \theta) = \frac{e^{\beta U(x, p(x), \theta)}}{\sum_{y \in X} e^{\beta U(y, p(y), \theta)}} \quad (4.2)$$

where $U(x, p, \theta)$ is some utility function smooth in p and $\beta \in [0, \infty)$ is the *attention parameter*.

In this chapter, I use AG price setting, which assumes that at every contract in X , there exists a fixed η -mass of “crazy types”—types of minimum cost or maximum value. η is an exogenous parameter, which we set by default to 10^{-8} , so as to not have any significant effect.

When considering informed buyers (e.g. insurance), we assume they have a cost $c(x, \theta)$, $dU/dp < 0$, and the price of each contract is

$$p(x) = \left[\int c(x, \theta) \mu(x, d\theta) + \eta \cdot \min_{\theta} c(x, \theta) \right] / (\mu(x, \Theta) + \eta)$$

Conversely, when considering informed sellers (e.g. job market), we assume they have a value $V(x, \theta)$, $dU/dp > 0$, and their price is

$$p(x) = \left[\int V(x, \theta) \mu(x, d\theta) + \eta \cdot \max_{\theta} V(x, \theta) \right] / (\mu(x, \Theta) + \eta)$$

4.3.3 Discussion

On the outflow. In linear logit, we completely abstract from passive choice to hone in on the active choice. Contracts are staggered a la Calvo (1983): agents change

their contract in any given period with a single exogenous probability α . In a model where the underlying distribution of θ is fixed over time, this could model birth-deaths, fixed-term contracts, and/or stochastic shocks to otherwise inertial agents, in each case forcing a new contract decision.

One technical drawback of the linear logit is that even with no inflow, the mass takes infinite time to fully decay (since $\mu' = -\alpha\mu$, $\mu(0) = 1$ is solved by $\mu = e^{-\alpha t}$). Other functional form specifications can guarantee that μ decays in finite time. Finally, linear logit falls into a class of models where outflow depends only on mass, meaning consumers are unaware of the emergence of particularly attractive options, in contrast to endogenous outflow models.

On (price) myopia. Unlike macroeconomic models, agents in DISE models are myopic insofar as they make decisions conditioned only on today's preferences. In theory, if utility is separable in x and p , the consumer could incorporate the fact that he is “stuck” with the contract terms (and the discounted utility flows arising from x) for a stochastic period of time, as is commonly done in the standard New Keynesian model, but nevertheless cannot forecast future changes in p . In many market settings, consumer myopia with respect to price is a reasonable assumption. Handel (2013), for example, argues in the context of health insurance:

...[prices] change as a function of factors that would be difficult for consumers to model. For consumers to understand the evolution of prices they would have to (i) have knowledge of the pricing model, (ii) have knowledge about who will choose which plans, and (iii) have knowledge about other employees' health.

Other types of economic models stipulate that one cannot “beat the market” by predicting consistently, so myopia holds by construction. Moreover, myopia can be a self-fulfilling mechanic of the market: myopia helps the model converge towards stable prices, and if prices are stable, that justifies price myopia.

The one situation which may be harder to fit within a myopic price model is one

in which prices can be expected, at least to go up or down. Those who choose to get a petroleum engineering degree, for example, are probably cognizant that the price of labour today does not reflect the price tomorrow. There are still ways to fold this into the model; for instance, we could encode into the contract space X , that by choosing to be a petroleum engineer, one must accept the disutility of certain consumption smoothing realities.

On multinomial logit. Logit is a sensible choice for a stochastic choice function—at least as a benchmark—for several reasons: i) it is transparent, analytically tractable, and computationally efficient ii) it is already a workhorse model in econometrics and iii) there has been a sizable literature on how logit can/cannot be justified (canonical examples include Luce (1959) and McFadden (2001)). In particular, Matějka and McKay (2015) showed that a weighted logit function arises when consumers are rationally inattentive with a cost of Shannon information inversely proportional to β . In their model, the weights are given by the prior probabilities, which helps the model overcome the irrelevant alternatives (also known as “red bus, blue bus”) critique of Debreu (1960).²⁷ Indeed, the model allows us to adjust for consumer biases for particular contracts by adjusting their prior probability of choosing that contract (for example, in insurance, there may be bias for zero coverage, maximum coverage, or some ‘default’). For simulation purposes, we simply define a suitable contract space X without close alternatives such that it is reasonable to assume a uniform prior.

Another interpretation of stochastic choice which is more appropriate when X represents a signal space is that signals are simply noisy. A law school applicant may

27. Debreu’s critique was that in logit models, the relative probability of taking a train vs. a red bus is equal to $\exp(\beta(U_{\text{train}} - U_{\text{red bus}}))$, and it remains so when half the red busses are replaced by blue ones. However, that cannot be, since the blue bus is a perfect substitute for the red bus. Matějka and McKay showed that if the rationally inattentive agent has a prior probability p_i over all forms of transportation $i \in T$, the stochastic choice function becomes a weighted logit $p_i \exp(\beta U_i) / \sum_j p_j \exp(\beta U_j)$, so probability ratio becomes $\frac{p_{\text{train}}}{p_{\text{red bus}}} \exp(\beta(U_{\text{train}} - U_{\text{red bus}}))$. When the prior probability of red busses are cut in half by the blue busses, the probability of taking the red bus halves. This weighting is equivalently thought of as shifting the utility by $\beta^{-1} \log p_i$.

intend to score a 170 on the Law School Admission Test, but may end up scoring 166 or 174. This is actually one of the more concrete examples in signaling—when signals are abstract messages open to interpretation by the market, it is difficult to reason how choice wouldn’t be stochastic, given a rich enough set of choices. Though there is not any reason *per se* why the stochastic choice function in this case would need to be multinomial logit.

On price setting. In competitive market models, when a contract is traded, price is equal to average cost. When the market is a selection market, however, the devil is in the detail of how prices are set when contracts are *untraded*, and specifically, how these shadow prices incentivize the unraveling of pools. Even more concretely, the simulated market should know to set the shadow price of a lower coverage contract low enough to incentivize good types to separate from bad.

The simplest mechanism involves setting the price of untraded contracts to the best possible price (the minimum cost or the maximum value). Under the right circumstances, this can draw away the types who would benefit from an inferior contract being offered at a lower price, if such types exist, and leave others behind. The untraded price holds for only one period, and may dissipate after trade occurs. This process of price discovery can be interpreted as entrepreneurs experimenting with the market by offering deals unprofitable in the short term, but after drawing away the right clientele, is sustainable in the longer term.

One clunky aspect of linear outflow models is that masses never fully decay, so we have to specify a condition under which a contract is considered “untraded”: say, if $\mu(x, \Theta) < \varepsilon$ for some $\varepsilon > 0$. To avoid erratic price setting behavior, one would have to specify $\alpha \ll \varepsilon \ll \frac{1}{|X|}$; this is in principle perfectly fine but in practice, the number of simulations required for convergence (in the colloquial, imprecise sense) is proportional to $1/\alpha$, meaning that such a price discovery procedure falls victim to a curse of dimensionality.

In simulations, however, nearly identical results are obtained by replacing the above process with AG price setting. The intuition is that, instead of assuming someone can cream skim by offering an untraded contract at the best price, we instead cheat by assuming that these contracts which are almost wholly untraded by regular types are nevertheless perpetually being traded at the best price to some tiny mass of crazy types, a mass which can be set arbitrarily small so that its influence on the nature of widely-traded contracts is imperceptible. AG price setting therefore emulates cream skimming, but allows us to remove a clunky ‘moving part’ in the model.

Framework flexibility. I would emphasize that the linear logit model is the *simplest* model within this DISE inflow-outflow framework, and a lot can be modelled simply by changing the nature of inflows and outflows. For example, perhaps passive choices also depend on the state of the market, and agents are more likely to outflow if the price of their contract becomes noticeably high compared to the market; there is nothing in the framework precluding the endogenization of outflows as well.

4.4 Job Market Signaling

4.4.1 Setup

We use the linear logit model from above. The signal space $x = 0, 0.1, 0.2, \dots, 10.0$ represents level of education. $\theta_1 = 1$, $\theta_2 = 2$, and $\theta_3 = 3$ each represent intrinsic productivity. Education is unproductive, thus the parameters of the model are

$$V(x, \theta) = \theta \qquad U(x, p, \theta) = p - \frac{x}{2\theta}$$

4.4.2 Theory

When $X = \mathbb{R}$ and Θ consists only of two types, both the D1 Criterion (Banks and Sobel 1987) and the Intuitive Criterion (Cho and Kreps 1987) eliminate all Perfect Bayesian Equilibria except the least-cost separating equilibrium, known as the *Ri-*

ley outcome (Riley 1979). In this setting, the Intuitive Criterion tends to be used because it tends to be regarded as more—for the lack of a better word—intuitive.²⁸ However, when Θ consists of three or more types, the Intuitive Criterion fails to eliminate all but the least-cost separating equilibrium, whereas the D1 Criterion does not fail. We use a type space with *three* types precisely to show whether and when the model converges to the D1 equilibrium, in part to provide a foundation for its usage.

To be specific, in the unique D1 equilibrium, all $\theta_1 = 1$ choose $x = 0$, all $\theta_2 = 2$ choose $x = 2$, and all $\theta_3 = 3$ choose $x = 6$. As shown in fig. 4.1, θ_1 chooses the smallest signal and receives wage $\theta_1 = 1$. θ_2 chooses the least costly signal capable of separating himself from θ_1 , by tracing along the indifference curve of θ_1 until it hits $p = \theta_2 = 2$. θ_3 separates from θ_2 analogously.

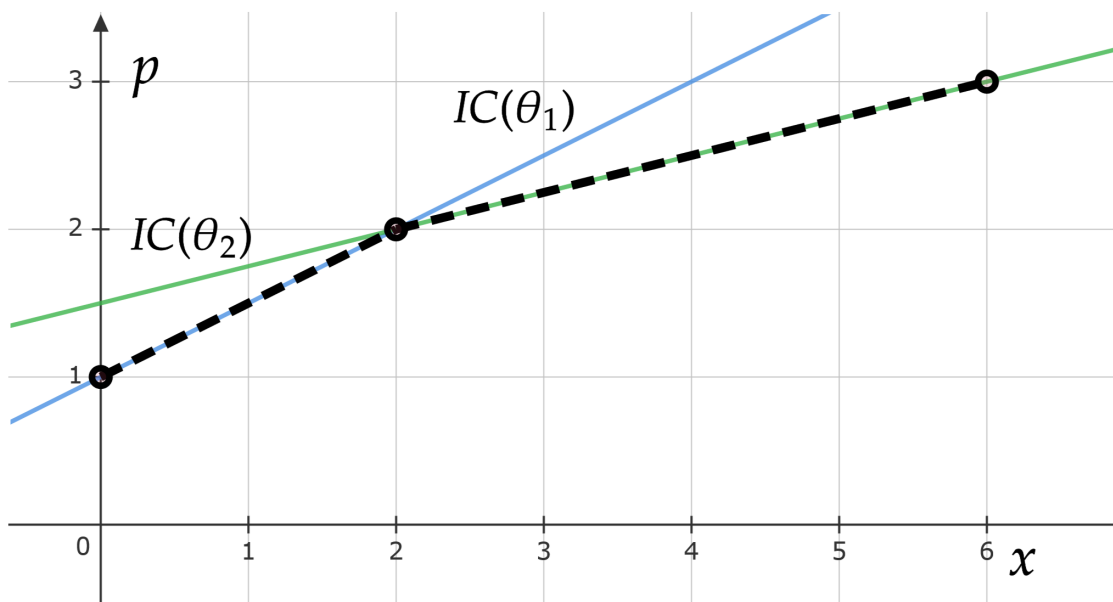


Figure 4.1: Riley outcome of Spence's signaling model

28. In particular, Cho and Kreps cite a “speech” that the signaling party ostensibly gives to convince the uninformed party not adopt unusual beliefs about off-equilibrium signals.

The D1, Intuitive, and Riley equilibria have all been subject to criticism (and even self-criticism, in the case of Azevedo and Gottlieb (2017), whose equilibrium produces a D1 result) on the basis that the outcome they produce are distribution free—that is, they depend only on the *support* of the distribution, and not on the distribution itself. The same equilibrium is produced when $P(\theta = 1) = \frac{1}{2}$ as when $P(\theta = 1) = 10^{-100}$. Mailath, Okuno-Fujiwara, and Postlewaite (1993) proposed an alternative in which those who separate consider the Riley equilibrium to which they would separate, and in particular, whether or not it “defeats” the pooling equilibrium. But this implies pooling in a two-type signaling case whenever the low type has probability $< \frac{1}{2}$, and this equilibrium seems more appropriate for strategic interactions than markets. In a dynamic competitive market, it doesn’t seem clear why high types couldn’t be gradually cream skimmed away from a pool even if they made up 51% of the pool, or why a separating equilibrium would collapse.

4.4.3 Results and Discussion

The Equidistributed Model. Assume that $\mu(\theta_1) = \mu(\theta_2) = \mu(\theta_3) = 1/3$. Unless otherwise specified, I set $\alpha = 0.001$, $\beta = 50$, the mass of AG behavioral types to 10^{-8} , and set the mass of all types at time zero to be $\mu^{t=0}(0, \Theta) = 1$; that is, all types start at no signaling.

The system settles at a D1 equilibrium, or at least distributionally close. Figure 4.2 shows the distribution of signals for $\beta = 50$ after 25,000 periods, starting with all types choosing zero signal.²⁹ It has three peaks, just north of $x = 0, 2, 6$, as D1 predicts. However, because choices are stochastic, there is some deviation from that prediction. This deviation is larger for θ_3 than for θ_1 . This makes sense from the rational inattention standpoint (as well as from a noisy signaling standpoint), since signaling is less expensive for θ_3 and thus over-signaling is less costly.

Figure 4.3 illustrates the convergence of the unconditional probability of each sig-

29. Not shown: I tried a variety of starting points, all of which converged at the same result.

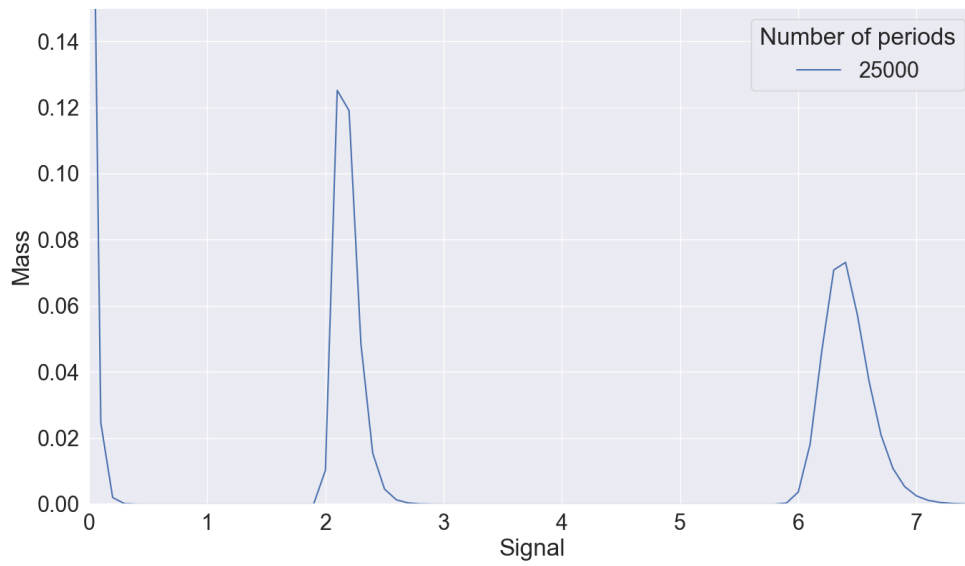


Figure 4.2: The “Noisy D1” in the equidistributed model

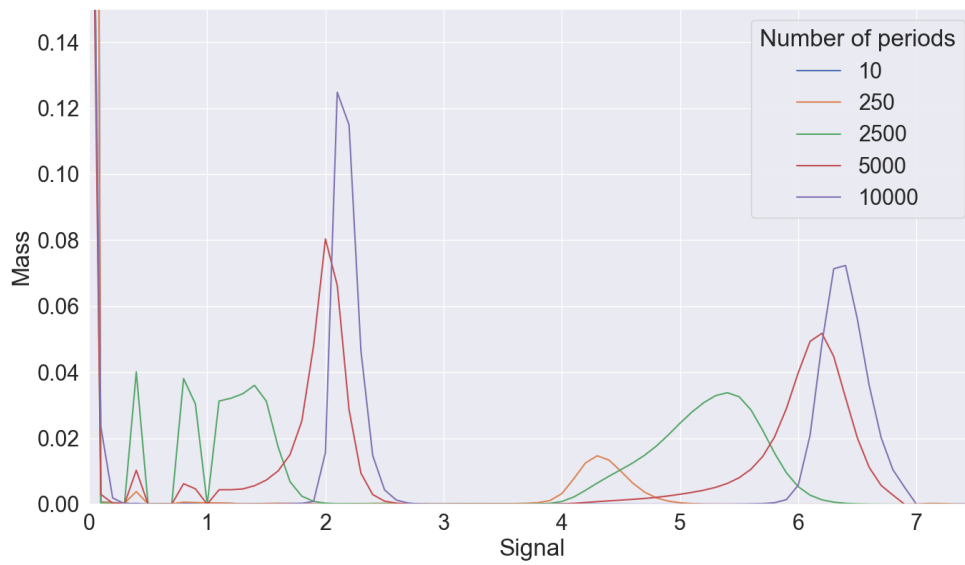


Figure 4.3: Convergence in the equidistributed model

nal $\mu^n(\cdot, \Theta)$ after n periods. We see that after roughly $n = 10^4$ periods, or $10 \times \alpha^{-1}$, the distribution has all but converged to the state shown in fig. 4.2.³⁰

Note that this convergence is not achieved if α is not sufficiently small and β is very large—instead, the system can behave chaotically. This is a natural consequence of inertia: if *everyone* were to move at once, then the optimal move in period n could become quite suboptimal by $n + 1$. So while it seems like inertia is an additional assumption added to the model, it actually alleviates us from making stringent assumptions about how these agents would behave strategically were they to act all at once.

As $\beta \downarrow 0$, *pooling becomes more pronounced*, as fig. 4.4 shows. When $\beta = 20$, the respective types are still largely separated. When β drops to 12, we still get a trimodal distribution, but one where there is significant “cross-subsidization” between θ_2 and θ_3 . Lowering β even further to 4, we obtain a bimodal distribution, with a local max in the middle of the signal space where marginals of θ_2 and θ_3 have converged to form a single peak, as the marginals show in fig. 4.5. Conversely, as $\beta \uparrow \infty$, there is less and less “noise” around the signals as more and more agents choose their D1 strategy.

Changing the distribution. I now turn to how the linear logit model deals with the canonical critique of Riley-like outcomes: namely, their distribution-freeness. Riley-like outcomes predict that, in this signaling model, there should be least-cost separation between the three different types, regardless of the proportion of each type, provided it is non-zero. A specific criticism is that it implies, regardless of how many θ_1 ’s there are, θ_2 separates from θ_1 (and likewise, θ_3 separates from θ_2).

30. When discussing simulations I use convergence colloquially here to mean that the distribution does not move discernibly after many periods; convergence (say in the sense of weak convergence of measures) is not proven, and as I discuss, for large α and β does not occur. Indeed, for general linear logit, one can construct simple counterexamples where the inferior type constantly tries to mimic the superior type, and the superior type constantly tries to separate, creating a “chase” through the contract space.

To explore this in the linear logit model, I set $\mu(\theta_1) = 1/r$ and $\mu(\theta_2) = \mu(\theta_3) = (1 - 1/r)/2$ for some $r > 1$. The other specifications remain unchanged from the equidistributed model.

As shown in fig. 4.6, when θ_1 types constitutes a relevant proportion of the population, θ_2 separates from θ_1 after 200,000 periods. Otherwise, it does not. The cutoff for “relevant” is somewhere in between $1/50$ and $1/100$.

What happens for some r in between 50 and 100 is that they converge to the D1 state (θ_2 separated from θ_1), but very slowly, and slower for larger r . What is shown in fig. 4.6 is that after 200,000 periods, the market with $\mu(\theta_1) = 1/89$ is stuck somewhere between the two equilibria. For some specifications of r , it may take millions of periods to converge.

Given that we only care about the short-to-medium run, it is interesting to inter-

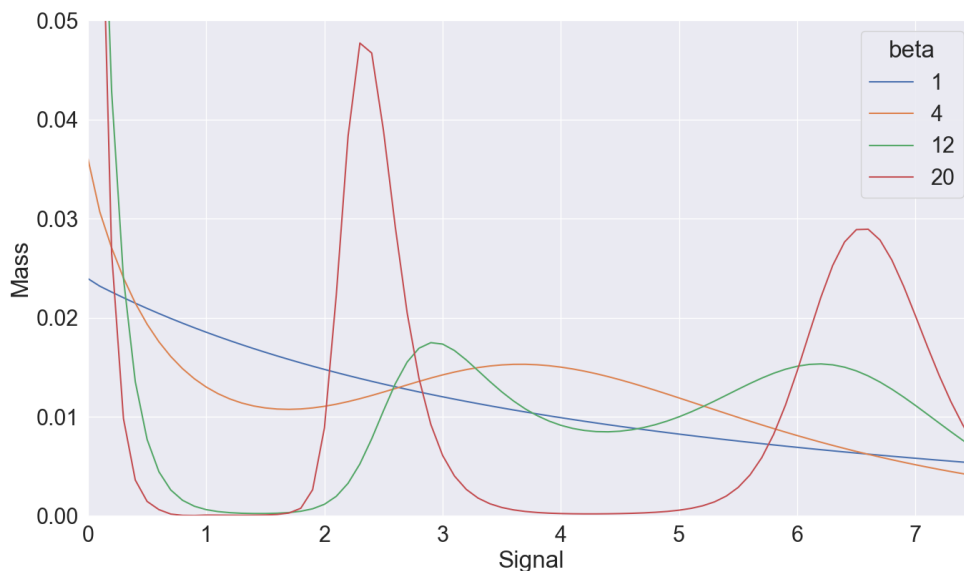


Figure 4.4: Limiting distributions as a function of β after 10,000 periods

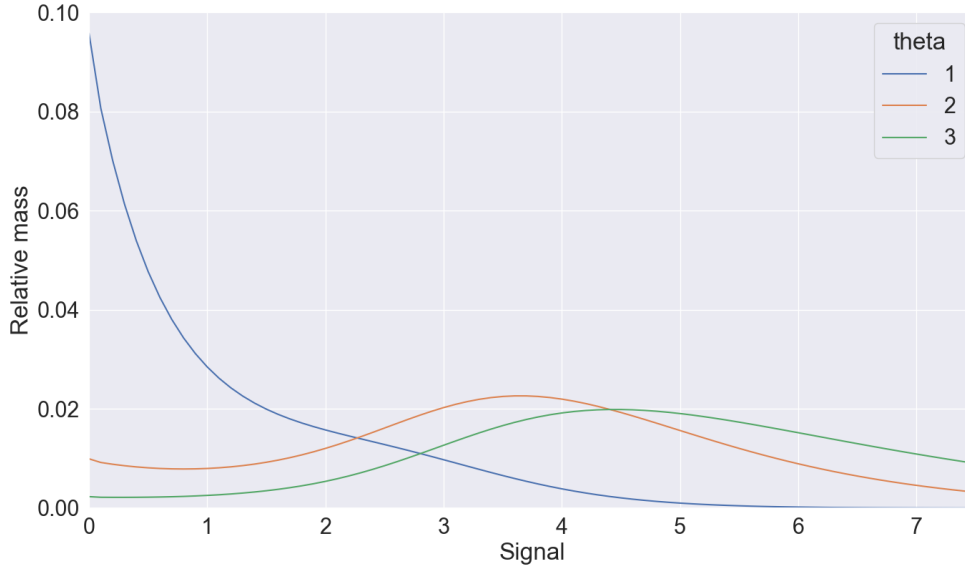


Figure 4.5: Marginal distributions for the case $\beta = 4$ after 25,000 periods

pret slow convergence as, in the short run, a “stable” state.³¹ This faux-equilibrium is actually quite intriguing, because we have θ_2 consistently choosing a signal between 0 and 2 (i.e. between its pooling signal and its D1 signal). The reason is that the 0 signal—because some θ_2 (and a small amount of θ_3 send it)—has an inflated wage relative to the fair wage for θ_1 , so it takes a smaller signal for θ_2 to separate, were he to act optimally.

Equilibria derived from theories of strategic interaction can destabilize in decentralized markets. I have briefly mentioned, to this point, that in the Spence signaling model (though not generally), the limiting distribution tends to not depend on the initial distribution.

But this leaves out a quite interesting story, which is how such an equilibrium

31. In the way that the Sun is pretty much constant in size throughout our lives, but will in billions of years expand to engulf what remains of our cold, dead bones.

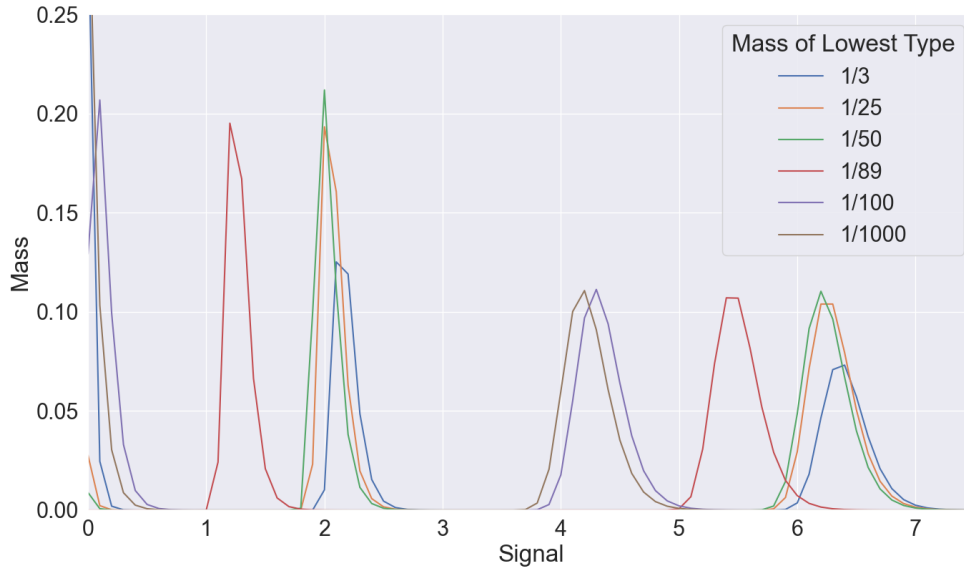


Figure 4.6: Signal distribution given $\mu(\theta_1)$ after 200,000 periods. All but 1/89 have converged.

is arrived. Figure 4.7 is particularly intriguing and illustrative. It shows that if we start with a sharp D1 equilibrium, the market does not converge nicely to a noisy D1. Rather, it collapses in on itself into a semi-separating equilibrium, with θ_2 's pooling with θ_1 's. Then, over millions of periods, the market cream skims the pool, slowly moving towards a separating equilibrium.

This is quite a clarifying example in that it highlights the difference between equilibrium concepts in strategic interactions—as characterized by equilibria in which “no one is better off deviating”—and de-centralized choices, even if only slightly stochastic, made in a market.

The vast majority of insurance buyers, for example, do not know how many people buy a particular contract, or what the least-cost separating equilibrium is, or what that term even means. They just respond to price. In the signaling game,

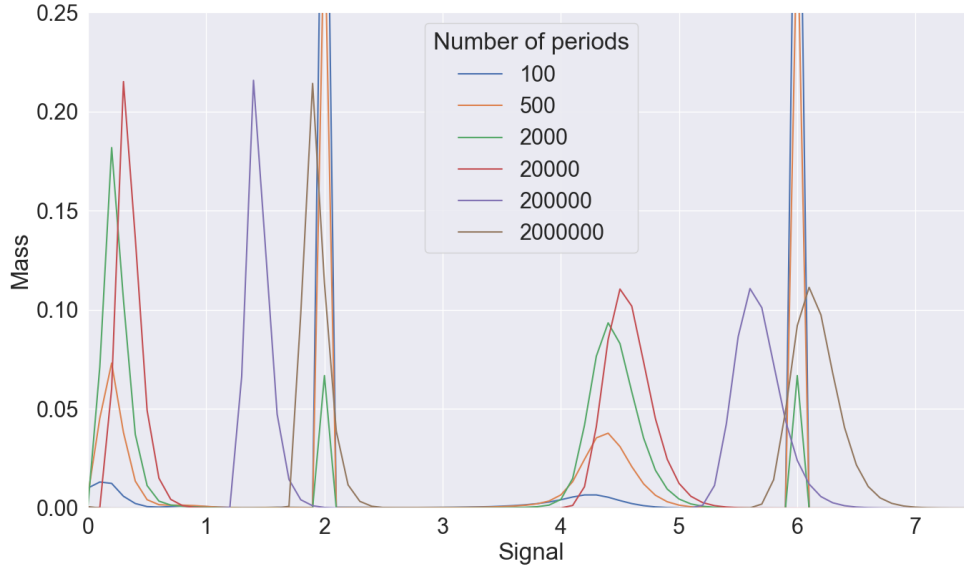


Figure 4.7: Convergence of the case $\mu(\theta_1) = 1/89$ when the initial distribution is the D1 equilibrium

the belief and the price are intertwined, so people do not “deviate” because they will be punished by the price. But in the DISE model, price is set by the movement and location of the agents in signal space. What happens after we set the initial condition at D1 is that θ_2 ’s, in stochastically choosing some smaller signals, quickly raise the average productivity associated with smaller signals. In doing so, they are “discovering” that they can get high wages sending small signals—and thus have no reason to send larger signals. It then takes time—in this case, a long time—for the market to slowly cream skim θ_2 ’s away from the θ_1 ’s by offering higher wages for higher signals.

4.4.4 Relationship to the Intuitive “Speeches”

One of the main critiques of D1 criterion and the Riley outcome at large is that it is unintuitive, at least in comparison to the Intuitive Criterion. The results suggest a new intuitive foundation for them: namely, that they represent what a stable

inflow-outflow equilibrium looks like in the limit, and more specifically, in a market where the uninformed side “experiments” with offering favourable prices on “off-the-equilibrium” contracts as a means of cream-skimming.

Adding noise to D1 also strengthens the intuitiveness of the proposed outcome, because another one of the critiques is that D1 puts too strong a condition on the out-of-equilibrium beliefs: the uninformed side of the market, in the single crossing case, assigns belief about out-of-equilibrium actions to only one player. In this market-based approach, beliefs about contracts which are dominated for all players (which one may term as “off-the-equilibrium”) are dictated by the empirical distribution of those who purchase it (except in the aforementioned “deals”), thus assigning belief to multiple players in a rational, Bayesian way.

The results are also intuitive in another sense. Recall that the “intuition” in Cho and Kreps (1987) is justified in part by the famed “speeches” that the informed player can give to eliminate “unintuitive” beliefs that the uninformed player may hold. The results also imply similar speeches that can be given in the market which, rather than eliminating “unintuitive” off-the-equilibrium beliefs, do the much simpler task of appealing to Bayesian plausibility (i.e. the “belief” of the “market” must be in accord with the empirical distribution of signals).

Consider the three-type (with low, medium, high types) signaling game with the following priors:

$$\begin{aligned}\mu(\theta_L) &= 0 \\ \mu(\theta_M) &= \varepsilon \\ \mu(\theta_H) &= 1 - \varepsilon\end{aligned}$$

for ε arbitrarily small. Once again, suppose productivity is given by $V(\theta) = \theta$ and the price of a contract (signal) is $p(x) = E[V(\theta)|x]$ given by the distribution of types who buys the contract (sends the signal).

The “intuitive” speeches, in this scenario, sound something like as follows. To eliminate any belief which would require θ_M to send any signal higher than 0, θ_M gives the following speech in justifying their action:

I am sending the signal 0, which ought to convince you that I am $> \theta_L$. For I would never (that is, with probability zero) both be θ_L and send signal 0, while if I am $> \theta_L$, and if sending this signal so convinces you, then, as you can see, it is in my interest to send it.

To justify separating from a pool by sending a profitable signal x_H , the least-cost separating signal between θ_M and θ_H assuming θ_M sends 0, θ_H recites a similar speech:

I am sending the signal x_H , which ought to convince you that I am θ_H . For I would never wish to send 0 if I was not θ_H , while if I am θ_H , and if sending this signal so convinces you, then, as you can see, it is in my interest to send it.

Thus, ruling out an M - H pooling equilibrium. A similar speech can also be made to justify sending a less-costly separating signal, thus ruling out any separating equilibrium except the least costly one.

Now compare this intuition with the market case. First, θ_M can choose signal 0 by making almost the exact same speech, but instead appealing to the fact that there are no θ_L 's, thus it would be non-Bayesian for the market to assign the signal 0 to θ_L 's.

Because $\mu(\theta_H)$ is arbitrarily close to 1, the competitive wage of a M - H pool will be arbitrarily close to θ_H . With a sufficiently fine signal space, θ_H can indeed make the speech and separate with a *slightly* higher signal; this changes the price and can eventually unravel the pool. However, suppose that the signal space is $X = \{x_1, \dots, x_n\}$ where $0 = x_1 < \dots < x_n$.³² The fact that X is finite (in the

32. for some sufficiently large $\bar{x}$ that the upper bound is not relevant in this analysis

DISE model and in reality) renders the task, in practice, implausible, at least in the short-and-medium term.

Consider a DISE model, except that there is one fully rational agent (me, without loss of generality).³³ First, consider the signals close to 0. For all x' close to 0, there will be some mass of θ_M 's, but θ_H will massively outnumber them because of stochastic choice. Thus,

$$p(0) \text{ is close to but less than } p(x') \text{ is close to but less than } \theta_H = V(\theta_H)$$

for any x' close to 0. Because the prices are so similar, there is almost zero incentive for θ_M to choose a slightly lower signal, or θ_H to choose a slightly higher one. In the rational inattention (logit) model with a uniform prior, the utilities being so similar implies the inflows are too (that is, because agents are nearly indifferent between signals, they basically choose at random). The speech of θ_H could sound something like this

I am sending some signal x' close to 0; regardless, you ought to assign a high probability that I am θ_H , since there is enough randomness in the choices of θ_H types that each signal near 0 is sent almost entirely by θ_H types.

Second, the discreteness of the signal X means that to separate, θ_H would need to send the second smallest signal x_2 , which may not be worth it at all, thus rendering the pool stable. For example, if the signal space near 0 consists of only $x_1 = 0$, and the utility for θ_H of receiving wage $V(\theta_H)$ sending x_2 is less than the utility of receiving wage $\int V d\mu(\theta)$ sending x_1 : in math, $U(x_2, \theta_H, \theta_H) < U(x_1, E\theta, \theta_H)$. The speech might sound something like this:

33. this one fully rational agent merely exists to analyze what a fully rational person *would* do in a world where most people act stochastically, and this person has zero mass

I am sending the signal $x_1 = 0$, which ought to convince you that I am probably θ_H . For most people who send signal 0 are θ_H , and no θ_H would want to send x_2 if the only incentive to do so is to go from convincing you that we're probably θ_H to convincing you that we're certainly θ_H .

But it is not only difficult to separate; for the same reason, separation is also easy to collapse. Consider an initial state which is the least-cost separating equilibrium (θ_M chooses 0, θ_H chooses x_H). At the initial state, the market's "belief," dictated by the empirical distribution of signals, is that the signal 0 is sent by θ_M and signal x_H is sent by θ_H .³⁴ To justify sending 0, I could just make the following speech:

I am sending the signal 0, which ought to convince you that I am probably θ_H , and thus deserve a wage close to $V(\theta_H) = \theta_H$. For so many θ_H will choose the signal 0 out of pure inattention in the next period that you will be forced raise my wage anyway. From then on, all θ_H will be incentivized, and increasingly so, to send the 0 signal instead of x_H . If sending this signal convinces you, then, as you can see, it is in my interest to send it.

4.5 Implications for Insurance Markets

4.5.1 Introduction: A "Model-Free" Analysis

Adverse selection causing a market to unravel provides a compelling rationale for government intervention. Some have alleged the problem is overblown and lacks empirical backing (Siegelman 2003), yet others have conjectured that the problem may be ubiquitous but difficult to observe since, by definition, an unraveled market is not there.

One recent empirical approach used to support the latter claim involves showing empirically that the private information in a market satisfies a theoretically derived

34. The market model does not have "off-the-equilibrium" beliefs *per se*, but we can suppose it entertains the D1 belief that $[0, x_H)$ is sent by θ_M and $[x_H, \infty]$ is sent by θ_H (on the other hand, under the linear logit model, the market prices all off-the-equilibrium contracts at θ_H , implicitly assigning all off the equilibrium contracts to θ_H)

no-trade condition, and thus precludes the existence of a market. Hendren (2013, 2017) centre on the following proposition:

Proposition 27 (Hendren (2013)). *Consider a model where types θ , drawn from distribution F , face a next-period loss l from their next-period wealth w with probability θ . They each have a von Neumann-Morgenstern utility function u . The market collapses – i.e. the only allocation which is resource feasible, individually rational, and incentive compatible is one with no insurance – if and only if*

$$\frac{u'(w-l)}{u'(w)} \leq \inf_{\tau \in \text{supp } F \setminus \{1\}} \frac{E[\theta|\theta \geq \tau]}{1 - E[\theta|\theta \geq \tau]} \frac{1-\tau}{\tau} =: T(\tau) \quad (4.3)$$

Furthermore, this no-trade condition holds only if there are perpetually worse types – that is, $1 \in \text{supp } F$.

The clever part of the proposition is that it is “model free” (in that it does not assume Akerlof-style fixed contracts or RS-style endogenous contracts). All it shows that, under the no-trade condition, one of the following pre-requisites of a stable insurance market must fail: resource feasibility, incentive compatibility, and individual rationality.

Hendren interprets the condition in detail in Hendren (2013). But basically, $\frac{E[\theta|\theta \geq \tau]}{1 - E[\theta|\theta \geq \tau]}$ is the fair cost of providing a contract which insures τ , assuming that any such contract will draw in all types $\theta \geq \tau$, while $\frac{u'(w-l)}{u'(w)} \frac{\tau}{1-\tau}$ is the willingness to pay of for an infinitesimal transfer of wealth from the event of a loss to the event of no loss. That is, a market exists if and only if some type τ , who knows that all risks worse than him will pool with him, is nevertheless willing to pay that price.

The results in the preceding section on the linear logit model suggest two effects could induce trade even when the no-trade condition holds.

The first effect, I regard as trivial: the inflow into any contract in the contract space is positive. This I do disregard as having little *per se* practical significance.

For one, if the model suggests a mass of, say, 10^{-7} , buying some contract, then that could only be regarded as an artifact of the DISE model without substantive economic relevance. But more importantly, in such a scenario, these types are not *willing* to buy the contract, they are merely making mistakes.

The second effect, though – that stochastic choice lowers the average cost of the contract with the riskiest pool – *is* significant, and may be compounded by the first. Under perfect rationality, trade occurs only if for some type τ is willing to pay the average cost when all types $\theta > \tau$ pool with him, and all types $\theta < \tau$ do not. But with stochastic choice, some of the riskier types $\theta > \tau$ will fail to buy into this pool and some types $\theta < \tau$ will accidentally buy into this pool. It follows that there are pathological cases where τ would not be willing to buy the contract if all buyers were rational, but *would* be willing provided sufficiently many agents made mistakes. I investigate this concept further in the context of the Akerlof and Rothschild-Stiglitz models.

4.5.2 The Akerlof Model

Typically, the fixed contracts paradigm works best with two insurance contracts, with high coverage H and low coverage L . For this section, assume that $H = \eta$ and $L = 0$, so that we may entertain the thought of an “infinitesimal” transfer of wealth from the state with loss to the state without.

Figure 4.8 illustrates the basic problem in selection markets.³⁵ Without adverse selection—for example, if people were to randomly, rather than selectively, buy insurance—one would expect the average cost to be a horizontal line AC_0 , since the average cost would involve an unconditional expectation over the risks in the population. In adverse selection markets, however, AC is downward sloping because higher price implies higher willingness to pay among buyers implies higher aver-

35. Note D and AC need not be linear; this is just an illustration

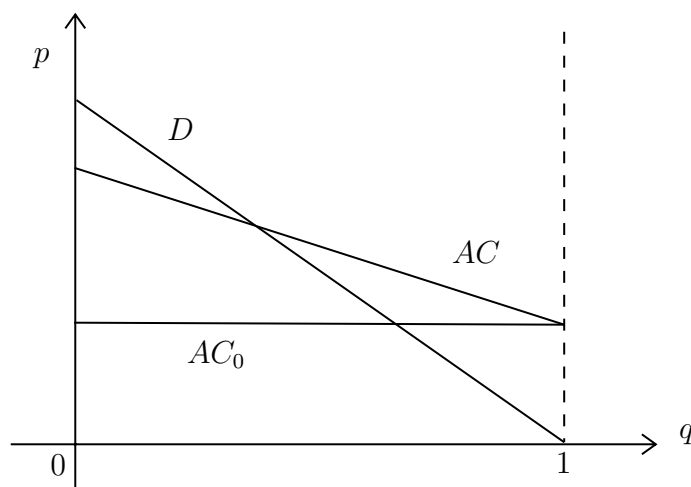


Figure 4.8: Akerlof unraveling of the market for high coverage contracts

age risk.³⁶ AC meets AC_0 only when $q = 1$ (here q describes the proportion of the population buying insurance), since that would return us to an unconditional expectation over the population. The demand curve D describes how many people would be willingness to pay at a certain price, and equilibrium arises where $D = AC$.

The complete unraveling case would be if the average cost curve AC laid strictly above the demand curve D , but even without complete unraveling, we see that adverse selection leads to an underprovision of insurance, compared to the case where we had constant costs AC_0 .

So what does inattention do? First, consider the average cost curve, and in particular, consider the extreme case where $\beta = 0$; this is just random choice, so the average cost curve would be AC_0 . Conversely, the case where $\beta = \infty$ corresponds to perfect rationality, and we return to the standard model where average cost is AC . So, it follows that inattention flattens the AC curve.

36. This is not always true beyond this textbook, one-dimensional case. For instance, higher willingness to pay could also imply greater risk aversion. In some cases, there can be *advantageous* selection. See Einav, Finkelstein, and Mahoney (2021).

Turning to the demand curve—in theory, willingness to pay stays the same, so the demand curve characterizing the mass q of people who are willing to pay at a given price p is unchanged. However, we could separately characterize, given a certain price, how many people would *actually* buy, call this the inattentive demand $\hat{D}$. Again, in the extreme case $\beta = \infty$, $\hat{D} = D$, and in the other case $\beta = 0$, $\hat{D}(p) = 0.5$; that is no matter the price, people pick randomly.³⁷ In between, $\hat{D}$ looks something like a logistic function: when price is very high, theoretically some tiny quantity still buys insurance, and vice versa. The graph of $\hat{D}$ steepens (flattens as a function of p) as β decreases (price of information increases) in that inattentive demand will respond less to the contract price as the information price increases.

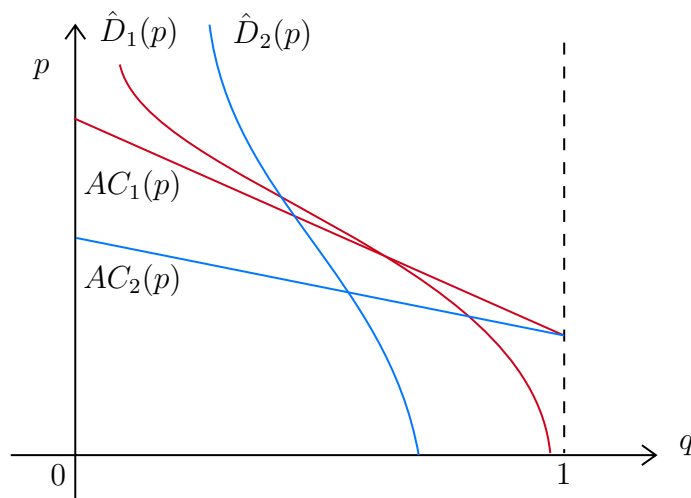


Figure 4.9: A change in β causing an ambiguous effect

Generally, the overall effect can be ambiguous, as shown in fig. 4.9. Holding $\hat{D}$ still, a flattening of AC induced by a decrease in β would increase quantity q and decrease price p . Likewise, when quantity is low and AC is held still, the steepening of $\hat{D}$ by decreasing β (which one could picture as pointwise convergence to $q = 0.5$, when $\beta = 0$) would increase quantity and reduce price. In either case, it's fairly plau-

37. in the Matějka-McKay model, one could shift the priors so that it's not 50-50, but we'll use 50-50 for now.

sible that welfare goes up even if some people aren't making the utility-maximizing choice. But if quantity is initially high (above the prior), the steepening of $\hat{D}$ while keeping AC still would decrease quantity and increase price.

Example. Here we illustrate a case where no type is willing to pay for an infinitesimal amount of insurance under perfect rationality, but not under stochastic choice. Consider the binary loss model with

$$u(c) = \log c \qquad l = \frac{w}{2} \qquad \theta \sim \text{Unif}[0, 1]$$

$E[\theta|\theta \geq \tau] = \frac{1-\tau}{2} + t = 1 - \frac{1-\tau}{2}$, and so $T(\tau) = 2$. On the other hand, $u'(c) = \frac{1}{c}$, so $\frac{u'(w-l)}{u'(w)} = \frac{w}{w-l} = 2$. Thus the no-trade condition (just barely) holds.

Now consider the market for an insurance contract which transfers an infinitesimal amount from the agent's no-loss state to his loss state. The willingness-to-pay (WTP_η) for a small finite transfer η satisfies

$$\begin{aligned} \theta \cdot u(w - l + \eta - WTP_\eta) + (1 - \theta) \cdot u(w - WTP_\eta) &= \theta \cdot u(w - l) + (1 - \theta) \cdot u(w) \\ &=: C, \end{aligned}$$

where C is a constant. The willingness-to-pay for a small amount of insurance normalized by dividing by the amount is approximated by its Taylor series, in turn given by the triple product rule:

$$\begin{aligned} \frac{WTP_\eta}{\eta} &\approx \left. \frac{dWTP_\eta}{d\eta} \right|_{\eta, WTP_\eta=0} = - \left. \frac{\partial C / \partial \eta}{\partial C / \partial WTP_\eta} \right|_{\eta, WTP_\eta=0} \\ &= \frac{\theta \cdot u'(w - l)}{\theta \cdot u'(w - l) + (1 - \theta) \cdot u'(w)}. \end{aligned}$$

It can be checked that the demand is

$$D(p/\eta) = P(WTP_\eta \geq p/\eta) = 2 - \frac{2}{2 - p},$$

while the average cost is

$$\frac{AC(q)}{\eta} = \eta^{-1} E[\theta \eta | \theta \geq 1 - q] = 1 - \frac{q}{2}.$$

The model is illustrated in fig. 4.10. As we see, the average cost curve lies above the demand curve, indicating full unraveling, just barely—which is what we expect since the no-trade condition holds, just barely.

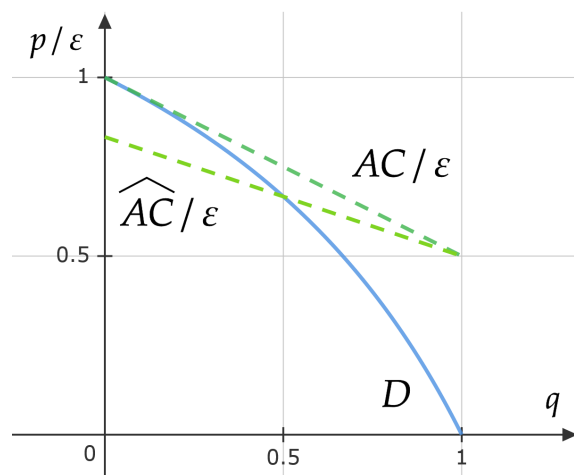


Figure 4.10: Example: Market for an infinitesimal amount of insurance

But we also see that if the average cost curve is lowered at all by stochastic choices, say to $\widehat{AC}$, then the curve *would* intersect the demand curve, indicating that there is a positive mass of types willing to pay for insurance.

4.5.3 The Riley-Rothschild-Stiglitz Model

As shown in Azevedo and Gottlieb (2017), a model with endogenous contracts yields an even starker result than the under-insurance result with fixed contracts: if $X = [0, 1]$ is the contract space and Θ is the support of the type distribution (describing the probability of binary loss), then the market unravels whenever $1 \in \Theta$. Taking into account fixed costs and moral hazard, the market unravels whenever

$\sup \Theta$ is sufficiently large.

The basic intuition is that, to create a least-cost separating equilibrium, types $\theta < 1$ would have to separate from $\theta = 1$. But that is not possible, because any insurance contract—which must be priced lower than the insured loss (that is, one would never buy pay a \$100 premium to insure against a loss of \$100)—would automatically draw in all types for whom $\theta = 1$.

The Willy Loman Principle. In Arthur Miller’s *Death of a Salesman*, Willy Loman, the salesman, pays premiums on a life insurance policy which he knows will be claimed. Would Willy Loman have single-handedly collapsed the life insurance market in New York?

The answer is obviously no, but it is harder to pinpoint exactly why. The stochastic choice answer is that, however separation occurs, the market does not separate others from Willy Loman—who, certain of his demise, buys the best contract—provided that Willy Lomans are vastly outnumbered by the mass of people who buy the same contract by chance, who are not certain of death.

Just as in the signaling case, if we simulate a case with two types, high risk θ_H and low risk θ_L , a pool which buys full coverage unravels quickly towards a (noisy version) least cost separation (no trade if $\theta_H = 1$) only if the mass of θ_H is sufficiently large. Otherwise, the pool unravels slowly or not at all.

4.6 A Health Insurance Model

With the previous section in mind as a simple model for why better decisions (by way of lower information costs) may *or* may not improve welfare, I show how we can use a DISE simulations to investigate counterfactuals.

Following Azevedo and Gottlieb (2017), this section attempts to illustrate the

equilibrium framework by constructing a model which, as faithfully as possible, replicates the model presented in Einav et al. (2013) (hereafter, EFRSC), and calibrates it to their parameter estimates.

4.6.1 Setup

The original model. First, I outline the setup in EFRSC. The market consists of privately informed buyers and zero-profit sellers. Buyers are defined by a multidimensional type $\theta = (\psi, \omega, \mu, \sigma^2)$. Buyers face a healthcare cost shock λ in period two, distributed according to a log-normal distribution with parameters (μ, σ^2) . Given their private information, they make a decision in period one to buy one of twenty-six evenly spaced contracts $X = \{0, 0.04, \dots, 1\}$, where x is their coverage against shocks. After receiving the shock, they choose a level of medical expenditure based on their shock λ and their *moral hazard type* ω ; skipping the derivations, their ex-post utility after optimal medical expenditures is

$$u(x, p, \theta; \lambda) = x\omega - \frac{(x\omega)^2}{2\omega} + y - (1 - x)\lambda - p$$

where y is income. Buyers have an ex-ante von Neumann-Morgenstern (vNM) utility function with constant absolute risk aversion (CARA)

$$U(x, p, \theta) = - \int \exp(-\psi u(x, p, \theta; \lambda)) dF_\theta(\lambda)$$

and the cost function is

$$c(x, \theta) = xE_\theta\lambda + x^2\omega$$

While I have formulated a linear logit model which differs from the model they estimated, I try to keep the linear logit model and its parameters as faithful to the original where possible. I discuss minor details of the model later on in this section.

For this exercise, α is set to 0.01, once more we consider a 26-contract space

$X = [0, 0.04, \dots, 1]$ with an initial state where all types choose full coverage, and run the model for 1000 periods.³⁸ For concreteness, this would model unraveling over the course of 1000 days, if 1% of the population in question were to make an active choice each day.

Finally, I would stress that this is purely an illustrative exercise to demonstrate the usefulness of the model in simulating counterfactuals. Inherent in the adoption of the EFRSC model is that we are co-opting the expected utility in EFRSC—used for ordinal purposes—for cardinal purposes in the logit function. Thus we make an implicit and arbitrary imposition upon the Matějka-McKay model: that the EFRSC model is true. This is, in part, to constrain the author’s degrees of freedom.

Details related to the simulation are found in section 4.8.2. Further discussion on the rational inattention foundations of the model, and some of the alternative assumptions that I could have used, are found in section 4.8.3.

4.6.2 Results

Distribution. The resulting distributions for $\beta = 10, 10^2, 10^4, 10^8, 10^{16}, 10^{32}$ are shown in fig. 4.11. Recalling that β^{-1} is the cost of information, the model suggests that lowering the cost (improving agents’ decisions) tends to lead to increased coverage in equilibrium, although the realized demand for full coverage is not necessarily monotonic. When information costs are small, the equilibrium is characterized in large part by indifference over high contracts.

Welfare. Taking the Matějka-McKay model as a microfoundation, there are two approaches to comparing welfare outcomes. The first is the expected contract utility:

$$W_C = \sum_{x \in X} \sum_{\theta \in \Theta} U(x, p(x), \theta) \cdot \phi(x, \theta)$$

38. after which we run for another 1000 to confirm that little has changed

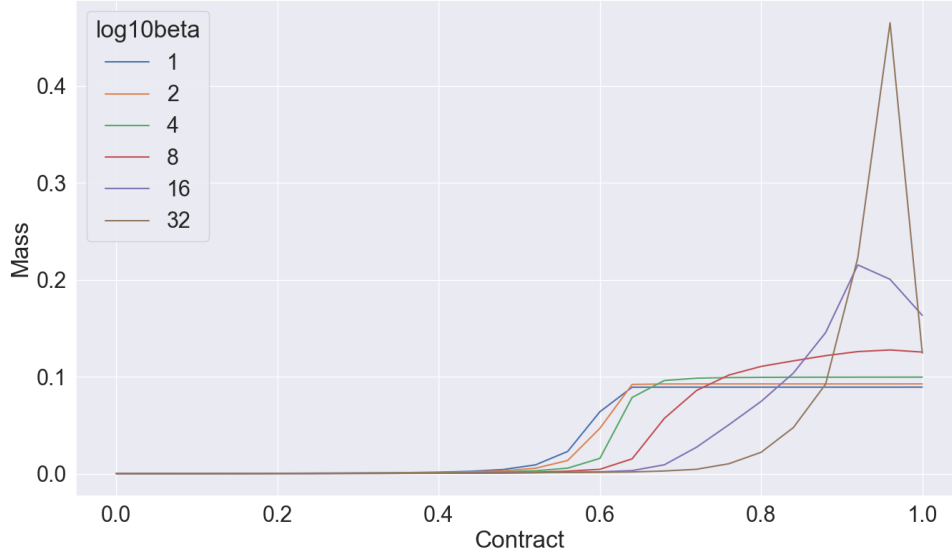


Figure 4.11: Distribution of coverage choices given various choices of $\log_{10}(\beta)$

which does not include the utility loss of information acquisition. The second, naturally, includes said loss:

$$W_T = -\beta^{-1} I_\phi(x; \theta) + \sum_{x \in X} \sum_{\theta \in \Theta} U(x, p(x), \theta) \cdot \phi(x, \theta)$$

Figure 4.12 shows the welfare effects of varying β , according to the model. Here too, the model suggests that lowering information costs improves overall welfare.

Caveats. It could be argued that one of the key features of the small β cases is actually an artifact of the CARA model. Namely, the fact that most types are indifferent some swathe of higher coverage contracts is in part a result of the upper boundedness of exponential utility; therefore, I would not make too much of the result. Even more specifically, because exponential utility is upper bounded, many agents will have some options which are far far away from the upper bound and some close to the upper bound; those far away will end up being assigned a probability near zero

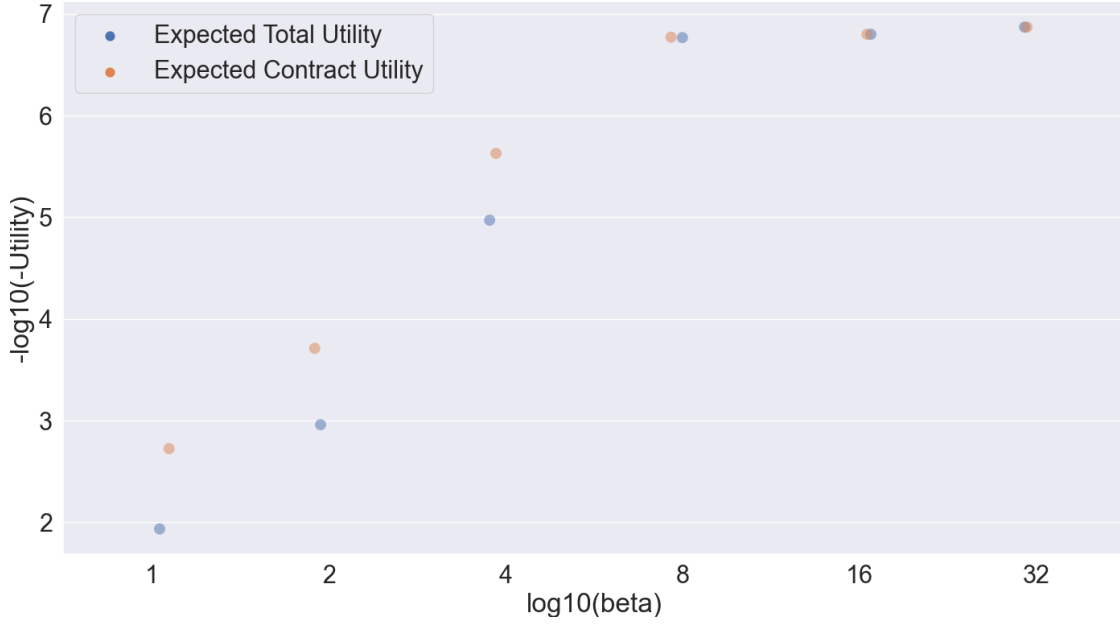


Figure 4.12: $-\log_{10}(-W)$ given six choices of $\log_{10}(\beta)$

by the softmax, and those close to 0 will be assigned a probability roughly equal to one another.

Another caveat relevant to small β cases is that baseline model assumes a uniform prior over all contracts, whereas typically, we might imagine more weight being put on large x and $x = 0$ (no insurance) and less on small positive values of x . Indeed, as the welfare results suggest, part of the reason that the distribution is more spread out when information costs are high is because information costs dominate in that scenario. When this is the case, the assumption regarding the priors have a more salient effect; what this also suggests is the importance of priors/choice architecture when information costs are high. Consider the extreme case where $\beta = 0$: in that case, the distribution of contract choices is determined *entirely* by the priors.

I see no negative implications of these caveats on the general efficacy of the DISE approach, as these problems can be easily fixed by calibrating the model to more

appropriately suit one's needs. I chose to follow a particular model specifically in the interest of constraining researcher degrees of freedom.

4.7 Concluding Remarks

The goal of this chapter was to trace out the implications of a fundamentally different set of assumptions about selection markets: instead of the traditional approach of identifying static equilibria with fully rational agents capable of changing contracts at any time, we study a market in which agents change contracts over time and, although they are irrational at a micro-level, their joint distributions are predictable at a macro-level.

To this end, this chapter provides a framework by which dynamic selection markets with choice frictions may be simulated. I construct a simple model under this framework to isolate the implications of rational inattention. In the classic one-dimensional case, the model produces a noisy least-cost separation which is far more robust than the equilibrium suggested by fully rational models. Finally I show that the model can be used to simulate potential counterfactuals in a calibrated model of a health insurance market.

4.8 Chapter Appendix

4.8.1 Models in the Limit

Even though the models don't work in continuous time, it is sometimes convenient to make use of calculus.

Define the total distance between the inflow and outflow by

$$G = \frac{1}{2} \sum_x \sum_{\theta} (\iota - \nu)^2$$

and its derivative with respect to time by

$$\frac{dG}{dt} = \sum_X \sum_{\Theta} (\iota - \nu) \left(\frac{d\iota}{dt} - \frac{d\nu}{dt} \right).$$

In the linear model, $\nu = \alpha(\iota - \nu)$. Holder's inequality gives

$$\begin{aligned} \frac{dG}{dt} &= -\alpha \sum_X \sum_{\Theta} (\iota - \nu)^2 + \sum_X \sum_{\Theta} 2(\iota - \nu) \left(\frac{d\iota}{dt} \right) \\ &= -2\alpha G + \sum_X \sum_{\Theta} (\iota - \nu) \left(\frac{d\iota}{dt} \right) \\ &\leq -2\alpha G + 2G \cdot \left\| \frac{d\iota}{dt} \frac{1}{\iota - \nu} \right\|_{\infty}. \end{aligned}$$

Thus, provided the norm is eventually bounded by α , $G \rightarrow 0$ eventually.

An interesting mathematical question, therefore, is whether there is always $\beta > 0$ which guarantees $\|d\iota/dt\|_{\infty}$ is bounded by β times a constant times $\|\iota - \nu\|_{\infty}$.

Proposition 28. *In the linear logit model, assume both $|U|$ and $\left| \frac{dU}{dp} \right|$ and bounded, if β satisfies*

$$\beta \left\| \frac{dU}{dp} \right\|_{\infty} [1 + (|X| - 1)e^{\beta(\bar{U} - U)}] \cdot \max_{x, \theta} \left| |\Theta| c(x, \theta) - \sum_{\theta' \in \Theta} c(x, \theta') \right| \cdot \max_{\theta} \mu(X, \theta) < 1,$$

then $G \rightarrow 0$ eventually.

Proof. β shows up in two places: first,

$$\begin{aligned} \frac{d\iota}{dt} \frac{1}{\iota} &= \frac{d\rho}{dt} \frac{1}{\rho} \\ &= \frac{1}{\rho} \left(\frac{d}{dt} \frac{e^{\beta U_x}}{\sum_X e^{\beta U_y}} \right) \\ &= \frac{1}{\rho} \left(\frac{e^{\beta U_x}}{\sum_X e^{\beta U_y}} \right) \cdot \beta \left(\frac{dU_x}{dt} - \frac{\sum_X \frac{dU_y}{dt} e^{\beta U_y}}{\sum_X e^{\beta U_y}} \right) \end{aligned}$$

$$\begin{aligned}
&= \beta \left(\frac{dU_x}{dt} - \sum_X \frac{dU_y}{dt} \rho_y \right) \\
&\leq \beta \max \left| \frac{dU}{dp} \right| \cdot \max \left| \frac{dp}{dt} \right|,
\end{aligned}$$

where the subscripts denote the relevant contract, when necessary.

Second, letting $(\bar{U} - \underline{U})$ be the maximum difference in utility between two contracts at price-equals-cost, the minimum inflow probability is

$$\begin{aligned}
\underline{\rho} &\geq \frac{e^{\beta \underline{U}}}{e^{\beta \underline{U}} + (|X| - 1)e^{\beta \bar{U}}} \\
&= \frac{1}{1 + (|X| - 1)e^{\beta(\bar{U} - \underline{U})}}.
\end{aligned}$$

The change in price is given by

$$\begin{aligned}
\frac{dp(x)}{dt} &= \frac{\sum_{\theta} c(x, \theta) \frac{d\mu(x, \theta)}{dt}}{\mu(x, \Theta)} - \frac{\sum c(x, \theta) \mu(x, \theta)}{\mu(x, \Theta)^2} \frac{d\mu(x, \Theta)}{dt} \\
&= \frac{1}{\mu(x, \Theta)} \left(\sum_{\theta} c(x, \theta) \frac{d\mu(x, \theta)}{dt} - p(x) \frac{d\mu(x, \Theta)}{dt} \right) \\
&= \frac{1}{\mu(x, \Theta)} \left(\sum_{\theta} [c(x, \theta) - p(x)] (\iota(x, \theta) - \nu(x, \theta)) \right).
\end{aligned}$$

Because inflow is always positive, we get that, in the long run, if $\iota(x, \Theta) \geq \underline{\iota}$, then

$$\begin{aligned}
\liminf_{t \rightarrow \infty} \mu(x, \Theta) &\geq \frac{\underline{\iota}(x, \Theta)}{\alpha} \\
&\geq \frac{\underline{\rho} \nu(X, \Theta)}{\alpha} \\
&= \underline{\rho}
\end{aligned}$$

and so,

$$\limsup_{t \rightarrow \infty} \frac{dp(x)}{dt} \leq [1 + (|X| - 1)e^{\beta(\bar{U} - U)}] \cdot \max_{\theta} \left| |\Theta|c(x, \theta) - \sum_{\theta' \in \Theta} c(x, \theta') \right| \cdot \|\iota - \nu\|_{\infty}.$$

Lastly, $\max_{\theta} \iota(x, \theta) \leq \max_{\theta} \nu(X, \theta) = \alpha \max_{\theta} \mu(X, \theta)$. Consequently,

$$\begin{aligned} \limsup_{t \rightarrow \infty} \left\| \frac{d\iota}{dt} \frac{1}{\iota - \nu} \right\|_{\infty} &= \limsup_{t \rightarrow \infty} \left\| \frac{d\iota}{dt} \frac{1}{\iota} \frac{\iota}{\iota - \nu} \right\|_{\infty} \\ &\leq \beta \left\| \frac{dU}{dp} \right\|_{\infty} [1 + (|X| - 1)e^{\beta(\bar{U} - U)}] \cdot \max_{x, \theta} \left| |\Theta|c(x, \theta) - \sum_{\theta' \in \Theta} c(x, \theta') \right| \cdot \alpha \max_{\theta} \mu(X, \theta). \end{aligned}$$

Lastly, dividing by α gives the desired result. $\square$

The bound here is, of course, crude. Several facts (and the discreteness of time in the actual model) suggest that heuristically, β need not be anywhere near the bound for the model to converge. First, a lot is lost by using Holder's inequality

$$\sum_X \sum_{\Theta} (\iota - \nu) \left(\frac{d\iota}{dt} \right) \leq 2G \cdot \left\| \frac{d\iota}{dt} \frac{1}{\iota - \nu} \right\|_{\infty}$$

since for the sum on the left, both $(\iota - \nu)$ and $(\frac{d\iota}{dt})$ are sometimes positive and sometimes negative, and tend to cancel out. Second, likewise, by examination of $\frac{dp}{dt}$, the sum which defines it is sometimes positive and sometimes negative.

4.8.2 Health Insurance Model Simulation Details

Following Azevedo and Gottlieb (2017), I consider 26 contracts specifying coverage level $x \in \{0, 0.04, \dots, 0.96, 1\}$. The types $\theta = (\psi, \omega, \mu, \sigma^2)$ lie on a $20 \times 20 \times 20 \times 12$ lattice, with the underlying probability mass function of each type corresponding roughly to the distributions accorded by the parameters estimated in EFRSC.

Computing the utility. The first trick we take advantage of is the fact that we can

decompose U into

$$\begin{aligned} U(x, p, \theta) &:= - \int \exp(-\psi u(x, p, \theta; \lambda)) dF_\theta(\lambda) \\ &= - \exp(\psi p) \underbrace{\exp\left(-\psi \left[x\omega - \frac{(x\omega)^2}{2\omega} + y\right]\right)}_{h(x; \theta)} \int \exp(\psi(1-x)\lambda) d\tilde{F}_\theta(\lambda). \end{aligned}$$

What this allows us to do is, instead of carrying out the full calculation each period, we compute h ahead of time and update the vNM via changes in price only.

In order to ensure that the integral defining U is $> -\infty$, instead of using a log-normal distribution F ,³⁹ we assume λ is drawn from $\tilde{F}$, a log-normal distribution truncated at $\bar{\lambda}$, which we set at \$75,000 based on the summary data in EFRSC.

While there isn't a clean formula for the moment-generating function of $\tilde{F}$, there is a closed form for its moments. Therefore, via Fubini's theorem, we change the integral to a sum:

$$\begin{aligned} \int \exp(\psi(1-x)\lambda) d\tilde{F}_\theta(\lambda) &= \sum_{n=0}^{\infty} \frac{\psi^n (1-x)^n}{n!} \int \lambda^n d\tilde{F}_\theta(\lambda) \\ &= \sum_{n=0}^{\infty} \frac{\psi^n (1-x)^n}{n!} \exp\left(\mu n + \frac{\sigma^2 n^2}{2}\right) \Phi\left(\frac{\log(\bar{\lambda}) - \mu}{\sigma} - \sigma n\right) \end{aligned}$$

As it turns out, we can bound the error of this approximation by the N -th term using Holder's inequality.

Lemma 29. *Shorthand the sum as $\sum_{n=0}^{\infty} \frac{\psi^n (1-x)^n}{n!} \int \lambda^n d\tilde{F}_\theta(\lambda) = \sum_n a_n$, we can bound the remaining terms by $\sum_{k=1}^{\infty} a_{N+k} \leq a_N \cdot \frac{\psi(1-x)\bar{\lambda}}{N+1-(\psi(1-x)\bar{\lambda})}$ whenever $N+1 > \psi(1-x)\bar{\lambda}$.*

^{39.} in EFRSC, they specifically actually assume $\lambda - \kappa$ is log-normal, where κ is yet another random variable. We omit this dimension.

Proof. We have

$$\begin{aligned}
a_{n+1} &= \frac{\psi^{n+1}(1-x)^{n+1}}{(n+1)!} \int \lambda^{n+1} d\tilde{F} = \frac{\psi^{n+1}(1-x)^{n+1}}{(n+1)!} \|\lambda^{n+1}\|_1 \\
&\leq \frac{\psi^{n+1}(1-x)^{n+1}}{(n+1)!} \|\lambda^n\|_1 \|\lambda\|_\infty = \frac{\psi(1-x) \|\lambda\|_\infty}{n+1} \frac{\gamma^n \psi^n (1-x)^n}{n!} \int \lambda^n d\tilde{F} \\
&= \frac{\psi(1-x)\bar{\lambda}}{n+1} a_n.
\end{aligned}$$

Shorthand $\zeta = \psi(1-x)\bar{\lambda}$. Then,

$$a_{N+k} \leq \frac{\zeta}{N+k} a_{N+k-1} = \frac{\zeta^2}{(N+k)(N+k-1)} a_{N+k-1} = \cdots = \zeta^k \frac{N!}{(N+k)!} a_N.$$

From whence we may obtain several representations of a bound on the sum of all terms greater than N whenever $N+1 > \zeta$

$$\begin{aligned}
\sum_{k=1}^{\infty} a_{N+k} &\leq a_N \sum_{k=1}^{\infty} \frac{\zeta^k N!}{(N+k)!} \\
&= a_N \left[\left(\sum_{k=0}^{\infty} \frac{\zeta^k N!}{(N+k)!} \right) - 1 \right] \\
&\leq a_N \left[\left(\sum_{k=0}^{\infty} \frac{\zeta^k}{(N+1)^k} \right) - 1 \right] \\
&= a_N \left(\frac{1}{1 - \frac{\zeta}{N+1}} - 1 \right) \\
&= a_N \cdot \frac{\zeta}{N+1 - \zeta}. \quad \square
\end{aligned}$$

The next lemma describes the actual summation we use to avoid over/underflow problems, since $\bar{\lambda}$ and $n!$ are naturally quite large.

Lemma 30. *We have*⁴⁰

$$h(x; \theta) = \frac{1}{2} \exp \left(\psi \left[\frac{x^2 \omega}{2} - x\omega - y \right] - \frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} \right) \\ \times \sum_{n=0}^{\infty} \exp \left(n \log(\psi(1-x)\bar{\lambda}) - \sum_{k=1}^n \log(k) \right) \cdot \xi_n$$

where $\xi_n = \operatorname{erfcx} \left(\frac{\sigma n}{\sqrt{2}} - \frac{\log(\bar{\lambda}) - \mu}{\sigma \sqrt{2}} \right)$.

Proof. We rewrite

$$\begin{aligned} & \Phi \left(\frac{\log(\bar{\lambda}) - \mu}{\sigma} - \sigma n \right) \\ &= \frac{1}{2} \exp \left(-\frac{1}{2} \left(\sigma n - \frac{\log(\bar{\lambda}) - \mu}{\sigma} \right)^2 \right) \operatorname{erfcx} \left(\frac{\sigma n}{\sqrt{2}} - \frac{\log(\bar{\lambda}) - \mu}{\sigma \sqrt{2}} \right) \\ &= \frac{1}{2} \exp \left(-\frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} + \sigma n \cdot \frac{\log(\bar{\lambda}) - \mu}{\sigma} - \frac{\sigma^2 n^2}{2} \right) \xi_n \\ &= \frac{1}{2} \exp \left(-\frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} + n (\log(\bar{\lambda}) - \mu) - \frac{\sigma^2 n^2}{2} \right) \xi_n \end{aligned}$$

and

$$\begin{aligned} & h(x; \theta) \\ &= e^{\psi(x^2 \omega / 2 - x\omega - y)} \sum_{n=0}^{\infty} \frac{\psi^n (1-x)^n}{n!} \int \lambda^n d\tilde{F} \\ &= e^{\psi(x^2 \omega / 2 - x\omega - y)} \sum_{n=0}^{\infty} \frac{\psi^n (1-x)^n}{n!} \frac{1}{2} \bar{\lambda}^n \xi_n \cdot \exp \left(-\frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} \right) \\ &= \frac{1}{2} \exp \left(\psi \left[\frac{x^2 \omega}{2} - x\omega - y \right] - \frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} \right) \sum_{n=0}^{\infty} \frac{\psi^n (1-x)^n \bar{\lambda}^n}{n!} \xi_n \end{aligned}$$

40. Note the abuse of notation in the representation of $n!$ as $\exp(\sum_{k=1}^n \log(k))$ which of course isn't valid for $n = 0$.

$$\begin{aligned}
&= \frac{1}{2} \exp \left(\psi \left[\frac{x^2 \omega}{2} - x\omega - y \right] - \frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} \right) \sum_{n=0}^{\infty} \exp \left(\log \left[\frac{\psi^n (1-x)^n \bar{\lambda}^n}{n!} \right] \right) \cdot \xi_n \\
&= \frac{1}{2} \exp \left(\psi \left[\frac{x^2 \omega}{2} - x\omega - y \right] - \frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} \right) \sum_{n=0}^{\infty} \exp \left(\log \left[\frac{\psi^n (1-x)^n \bar{\lambda}^n}{n!} \right] \right) \cdot \xi_n \\
&= \frac{1}{2} \exp \left(\psi \left[\frac{x^2 \omega}{2} - x\omega - y \right] - \frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} \right) \\
&\quad \times \sum_{n=0}^{\infty} \exp \left(n \log(\psi(1-x)\bar{\lambda}) - \sum_{k=1}^n \log(k) \right) \cdot \xi_n. \quad \square
\end{aligned}$$

For the purposes of our simulation, we used the bound

$$\begin{aligned}
10^{-14} &\geq \exp \left(\frac{(\frac{\log(\bar{\lambda}) - \mu}{\sigma})^2}{2} \right) \sum_{n>N}^{\infty} \frac{\psi^n (1-x)^n}{n!} \int \lambda^n d\tilde{F}_{\theta}(\lambda) \\
&= \sum_{n>N}^{\infty} \exp \left(n \log(\psi(1-x)\bar{\lambda}) - \sum_{k=1}^n \log(k) \right) \cdot \xi_n
\end{aligned}$$

4.8.3 Rational Inattention Foundations of the Utility Function

Consider a rational inattention model in which agents are not fully aware of their type θ – say, they do not know μ or σ^2 , but definitely know ψ and ω . Before the game, they commit to an information strategy described by a joint distribution between optimal contract and type $\phi(x, \theta)$ over the probability space $(\Omega, \mathcal{F}, P)$. In the first period, they receive a signal and choose a contract based on the information strategy. In the second, the true θ is revealed to them. In the third period, they receive the health shock and take the optimal health expenditure action in the third period. Matějka and McKay (2015) show that any agent solving an expected-utility-minus-

information-cost problem

$$\begin{aligned}
& \max_{\phi} -\beta^{-1} I_{\phi}(x; \theta) + \sum_{x \in X} \sum_{\theta \in \Theta} U(x, p, \theta) \cdot \phi(x, \theta) \\
& \text{subject to Bayes plausibility: } \sum_{x \in X} \phi(x, \theta) = P(\theta) \\
& \text{where } I_{\phi}(x, \theta) = \sum_{x \in X} \sum_{\theta \in \Theta} \phi(x, \theta) \log \left(\frac{\phi(x, \theta)}{P(x)P(\theta)} \right)
\end{aligned}$$

where $P(\theta)$ is the prior of θ , $I(x; \theta)$ is the Shannon mutual information, and β^{-1} is the search cost, the logit stochastic choice function uniquely solves the optimization problem.

The astute reader will notice a slight but not unresolvable hitch with the Matějka-McKay story when applied here: rational inattention is all about the agent not knowing θ , whereas selection markets are all about the agent being the only one who knows θ . The resolution I propose is as follows. Suppose $\theta = (\theta^k, \theta^a)$ representing the information which is *known* to the agent and the information which must be *acquired* by the agent. Regardless of whether θ , or just θ^a , represents information which must be acquired, it will nevertheless be the case that

$$P(x|\theta^k, \theta^a) = \frac{e^{\beta U(x, p(x), \theta^k, \theta^a)}}{\sum_y e^{\beta U(y, p(y), \theta^k, \theta^a)}} = \frac{e^{\beta U(x, p(x), \theta)}}{\sum_y e^{\beta U(y, p(y), \theta)}} = P(x|\theta)$$

So while from the view of the agent, there is a difference between not knowing θ and not knowing only θ^a , from the point of view of the market which does not observe any part of θ , $P(x|\theta)$ remains multinomial logit.

The remaining discussion to come centers on the unenviable task what cardinal U to use. In the health insurance model, I take U simply to be the expected CARA utility (over health states) in EFRSC, which is ordinal; in the Matějka-McKay model, U is cardinal, as the model involves taking an expectations over (x, θ) . Any vNM utility function U which is a strictly increasing transformation of

$-\int \exp(-\psi u(x, p, \omega; \lambda)) dF_\theta(\lambda)$ will maintain the preference over choices $x \in X$ conditional on type. None inherently make more sense than the one I used, but to give an idea of the other possibilities, I consider some of the implications of the following candidates for U :

1. the “CARA”: $U_{CARA} = -\int \exp(-\psi u(x, p, \omega; \lambda)) dF_\theta(\lambda)$
2. the “adjusted CARA”: $U_{adj} = -\gamma^{-1} \int \exp(-\psi \gamma u(x, p, \omega; \lambda)) dF_\theta(\lambda)$
3. the “ L^p ”: $U_L = -\gamma^{-1} \left(\int \exp(-\gamma u(x, p, \omega; \lambda))^\psi dF_\theta(\lambda) \right)^{1/\psi}$
4. the “log L^p ”: $U_{LL} = -\frac{1}{\psi \gamma} \log \left(\int \exp(-\psi u(x, p, \omega; \lambda)) dF_\theta(\lambda) \right)$

Interpretation. The CARA utility, when plugged into the maximization problem, the resulting expected utility can be interpreted as an application of the tower property:

$$\begin{aligned}
& \sum_{x \in X} \sum_{\theta \in \Theta} - \int \exp(-\psi u(x, p, \omega; \lambda)) dF_\theta(\lambda) \cdot \phi(x, \theta) \\
&= E_\phi [E[-\exp(-\psi u(x, p, \omega; \lambda)) | \theta, x]] \\
&= E_\phi [-\exp(-\psi u(x, p, \omega; \lambda))] .
\end{aligned}$$

On the surface, this makes the most sense, but this actually ignores an important aspect of the structure of the model. The agent is not choosing a contract conditional on θ , rather, the agent chooses, and importantly, *commits* to an information strategy.

To make things more concrete, consider someone choosing an insurance policy x , who can choose how much research to do into the state of the world θ . Maćkowiak, Matějka, and Wiederholt (2023) notes that, similar in spirit to Kamenica and Gentzkow (2011), these two choices are tantamount to a single commitment to a Bayes plausible joint distribution $\phi(x, \theta)$. It is only after choosing the policy that the state of the world θ is revealed – for example, the onset of a particular medical condition. No health expenditure shock has yet hit, but he has an interim utility U —a certain “(un-)happiness” about being stuck with (x, p) now that θ is known—which is a strictly increasing function of $E[-\exp(-\psi u(x, p, \omega; \lambda))]$, but need not be

a linear function.⁴¹

L^p gives a particular way for the agent to evaluate his contentment with a number of scenarios (x, θ) . Consider what the agent with $\psi = \infty$ (i.e. the sup norm) is doing in the rational inattention paradigm. He sits at the kitchen table evaluating policies over states of the world and, to each (x, p, θ) , assigns a utility $U(x, p, \theta) = \text{ess sup}_\lambda e^{-\psi \gamma u(x, p, \omega; \lambda)}$, corresponding to essentially⁴² the worst-case health outcome in that state of the world. Conversely, as $\psi \rightarrow -\infty$, he assumes best case scenario. In this view, ψ can be seen as incorporating “optimism” or “pessimism”. Insofar as the usefulness of insurance is partially derived from providing peace-of-mind, this provides for an interesting interpretation of the parameters.

Price levels. It is easy to see that for any constant η ,

$$u(x, \eta p, \eta \omega; \eta \lambda) = \eta u(x, p, \omega; \lambda).$$

This poses a theoretical problem for CARA which is solved by the adjustment. Namely, the odds ratio of two contracts is dependent on price level. Recalling that β^{-1} is the price of information:

$$\frac{P(x_1|\theta)}{P(x_2|\theta)} = \exp \left(\beta \eta^{-1} (E[-\exp(-\psi \eta u_1)] - E[-\exp(-\psi \eta u_2)]) \right).$$

Rates and elasticities of substitution. By assumption, $u(x, p, \omega; \lambda)$ is quasilinear, so we can move $w = y - p$ outside of U . However, in each case, a different thing is moved out:

$$U_{adj} = -\gamma^{-1} \exp(-\psi \gamma w) \int \exp(-\psi \gamma u(x, \omega; \lambda)) dF_\theta(\lambda)$$

41. In the end, the “expected utility” in fig. 4.12 represents the preferences of agents from an ex-ante perspective, in the sense of a Rawlsian veil of ignorance: expected utility over all states of the world without any information about it, though we may make a conditional expected utility calculation by holding some dimension(s) of θ fixed and integrating over the remaining ones.

42. in the measure-theoretic sense and in the colloquial sense

$$U_L = -\gamma^{-1} \exp(-\gamma w) \left(\int \exp(-\gamma u(x, \omega; \lambda))^\psi dF_\theta(\lambda) \right)^{1/\psi}$$

$$U_{LL} = w - \frac{\gamma}{\psi} \log \left(\int \exp(-\psi \gamma u(x, \omega; \lambda)) dF_\theta(\lambda) \right).$$

Given a fixed ϕ such that

$$V^* = -\beta^{-1} I_\phi(x; \theta) + \sum_{x \in X} \sum_{\theta \in \Theta} U(x, p, \theta) \cdot \phi(x, \theta)$$

the marginal rates of substitution between information acquisition cost and numeraire wealth loss, denoted $MRS = \frac{\partial V / \partial(-w)}{\partial V / \partial I}$ are

$$MRS_{adj} = \left(\frac{\psi \gamma e^{-\psi \gamma w}}{\beta^{-1}} \right) \sum_{x \in X} \sum_{\theta \in \Theta} \gamma^{-1} \int \exp(-\psi \gamma u(x, \omega; \lambda)) dF_\theta(\lambda) \cdot \phi(x, \theta),$$

$$MRS_L = \left(\frac{\gamma e^{-\gamma w}}{\beta^{-1}} \right) \sum_{x \in X} \sum_{\theta \in \Theta} \gamma^{-1} \left(\int \exp(-\gamma u(x, \omega; \lambda))^\psi dF_\theta(\lambda) \right)^{1/\psi} \cdot \phi(x, \theta),$$

$$MRS_{LL} = \beta.$$

If we interpret “cost” β of information literally, then naturally, $\log L^p$ is the most reasonable, since the tradeoff is linear.

Now assume that I is not numeraire, but rather comes about via costly search (loss of leisure as opposed to a loss of money). Take $s = e^{\xi I}$ to be the search cost. Then, letting C be a constant such that $MRS_L(I, -w) = C e^{-\gamma w}$:

$$MRS_L(s, -w) = MRS_L(I, -w) \frac{ds}{dI} = C e^{-\gamma w} \xi s$$

The derivative of the MRS with respect to $s/(-w)$ is

$$\frac{dMRS_L(s, -w)}{d(-s/w)} = - \frac{dMRS_L(s, -w)}{ds} \frac{ds}{d(-s/w)} = -C w e^{-\gamma w} \xi$$

Thus the nice thing about CARA/L^p is that the elasticity of substitution is 1:

$$\frac{MRS_L(s, -w)}{-s/w} \bigg/ \frac{dMRS_L(s, -w)}{d(-s/w)} = -\frac{Ce^{-\gamma w} \xi s w}{s} \bigg/ -Cwe^{-\gamma w} \xi = 1$$

Non-stochastic implications (F_θ is Dirac). Thus far we have remained within the rational inattention framework. Here we ask, given our functional form for U , what is the implied probability that someone “makes a mistake” when choosing an insurance policy (for example, by choosing a policy which is strictly dominated by an alternative)?

To fix ideas, perhaps it is better to ask, what is the probability that someone fails to pick a \$10 dollar bill off the ground if the bill is there almost surely? For adjusted CARA, $U = -\gamma^{-1} \int \exp(\gamma \psi \chi m) dP(m) = -\gamma^{-1} \exp(\gamma \psi \chi m)$, where χ is the indicator for whether the bill is picked up and P is a Dirac probability with $P(m = 10) = 1$. The ratio between the probability of picking up a bill on the ground vs. not implied by the rational inattention model is $\frac{P(\chi=1|\theta)}{P(\chi=0|\theta)} \exp(\beta[U(1) - U(0)]) = \exp(\beta\gamma^{-1}[-\exp(-10\gamma\psi) + 1])$, meaning that risk averse agents *are* more likely to pick up the bill, and this is true *irrespective* of the information cost β^{-1} . For L^p , $U = -\gamma^{-1} \exp(\gamma \chi m)$, where γ is a constant, so the likelihood ratio is independent of ψ , so risk aversion plays no role. For log L^p , $U = -\chi m$, which is likewise independent of ψ but also independent of the price level γ .

Again, there is no right way to interpret this *per se* – it depends on the mental model that one has about insurance choice. On the one hand, it seems plausible in the context of rational inattention that more risk averse people spend longer around the kitchen table thinking about what the future state of the world might be. On the other, if we interpret this outside the context of rational inattention, it seems plausible that simple mental lapses might be orthogonal to risk aversion.

Chapter 5

Conclusion

This thesis provides several new results in the theory and simulation of selection markets. The first reconciles the Riley and Azevedo-Gottlieb equilibrium concepts by laying out the conditions under which the Riley IVP provides a unique SAGE. The second augments the standard Riley model to include a fixed cost and provides a condition under which the fixed cost fully unravels the market. The final chapter sets out a broader framework for simulating selection markets, and uses it to study the implications of rational inattention on selection markets, finding that inattention adds noise to the standard equilibrium, making equilibria more stable and less prone to unraveling.

Future theoretical work on selection markets can contribute in a number of directions: i) to the way we understand economic forces, ii) to the way we understand existing models, and iii) to our ability to make predictions.

On the first point, very little substantively has changed on this front in recent decades. Though models have been produced to better simulate multidimensional models and (herein) models with choice frictions, there is not yet a theoretical framework to think about multidimensional or dynamic problems on par with the way we think about unidimensional static problems in terms of Akerlof, Spence, and RS. That is, we have a much easier time explaining why there might be a correlation be-

tween risk and coverage than explaining why there might not be a correlation when risk aversion comes into play, as in Finkelstein and McGarry (2006). Future work could therefore improve the way we think about such markets.

On the second point, Azevedo-Gottlieb equilibrium is arguably the first major advance in the theory of competitive selection markets in a number of years, and that it essentially replicates the original reactive equilibrium adds to its credibility. However, it remains unwieldy to use as a pen-and-paper tool. Future work could aim to make it more tractable, or connect it more broadly to qualitative ideas about selection market behavior, or replace it altogether with a new approach.

On the third point, models which produce least-cost separating equilibria in the single-crossing case still face the challenge of empirical testing. In the spirit of Hendren (2013, 2017), the second chapter’s main result could be empirically tested by looking at whether market non-existence may be predicted by the existence of high-risk types not willing to pay their fair price for insurance.

Lastly, dynamic models of selection markets open up a large array of theoretical questions. One direction could be to trace out the implications of other behavioral models. Another could be to build a model incorporating more behavioral elements to make quantitatively better predictions about the trajectory of selection markets. A final question, opened up by our exploration of the effect of information costs on unraveling, is a question of designing the marketplace for inattentive consumers: a designer (e.g. a government in an insurance marketplace, or the human resources department in a workplace insurance scheme) could save buyers from information acquisition costs by giving good recommendations (i.e. the designer commits to providing a recommendation that is close enough to optimal for the buyer that they won’t be incentivized to acquire much more costly information); in doing so, it has some control over the outcome of the insurance market. This opens up new research directions into what an optimal information policy could be.

Bibliography

- Abaluck, Jason, and Jonathan Gruber. 2011. “Choice inconsistencies among the elderly: Evidence from plan choice in the Medicare Part D program.” *American Economic Review* 101 (4): 1180–1210.
- Akerlof, George A. 1970. “The market for “lemons”: Quality uncertainty and the market mechanism.” *The Quarterly Journal of Economics* 84 (3): 488–500.
- Azevedo, Eduardo M, and Daniel Gottlieb. 2017. “Perfect competition in markets with adverse selection.” *Econometrica: Journal of the Econometric Society* 85 (1): 67–105.
- . 2019. “An example of non-existence of Riley equilibrium in markets with adverse selection.” *Games and Economic Behavior* 116:152–157.
- Banks, Jeffrey S, and Joel Sobel. 1987. “Equilibrium selection in signaling games.” *Econometrica: Journal of the Econometric Society* 55 (3): 647–661.
- Boehm, Eduard. 2024. “Intermediation, choice frictions, and adverse selection: Evidence from the Chilean pension market.”
- Brown, Zach Y, and Jihye Jeon. 2023. “Endogenous information and simplifying insurance choice.”
- Calvo, Guillermo A. 1983. “Staggered prices in a utility-maximizing framework.” *Journal of Monetary Economics* 12 (3): 383–398.

- Chandra, Amitabh, Benjamin Handel, and Joshua Schwartzstein. 2019. “Behavioral economics and health-care markets.” In *Handbook of Behavioral Economics: Applications and Foundations 1*, 2:459–502. Elsevier.
- Chen, Chia-Hui, Junichiro Ishida, and Wing Suen. 2022. “Signaling under double-crossing preferences.” *Econometrica: Journal of the Econometric Society* 90 (3): 1225–1260.
- Chetty, Raj, and Amy Finkelstein. 2013. “Social insurance: Connecting theory to data.” In *Handbook of Public Economics*, 5:111–193. Elsevier.
- Chiappori, Pierre-André, and Bernard Salanié. 2000. “Testing for asymmetric information in insurance markets.” *Journal of Political Economy* 108 (1): 56–78.
- Cho, In-Koo, and David M Kreps. 1987. “Signaling games and stable equilibria.” *The Quarterly Journal of Economics* 102 (2): 179–221.
- Cho, In-Koo, and Joel Sobel. 1990. “Strategic stability and uniqueness in signaling games.” *Journal of Economic Theory* 50 (2): 381–413.
- Cutler, David M, Amy Finkelstein, and Kathleen McGarry. 2008. “Preference heterogeneity and insurance markets: Explaining a puzzle of insurance.” *American Economic Review* 98 (2): 157–162.
- Debreu, Gerard. 1960. “Review of ‘Individual choice behavior: A theoretical analysis’ by R. D. Luce.” *The American Economic Review* 50 (1): 186–188.
- Einav, Liran, Amy Finkelstein, and Neale Mahoney. 2021. “The IO of selection markets.” In *Handbook of Industrial Organization*, 5:389–426. 1. Elsevier.
- Einav, Liran, Amy Finkelstein, Stephen P Ryan, Paul Schrimpf, and Mark R Cullen. 2013. “Selection on moral hazard in health insurance.” *American Economic Review* 103 (1): 178–219.
- Engers, Maxim, and Luis Fernandez. 1987. “Market equilibrium with hidden knowledge and self-selection.” *Econometrica: Journal of the Econometric Society* 55 (2): 425–439.

- Finkelstein, Amy, and Kathleen McGarry. 2006. "Multiple dimensions of private information: Evidence from the long-term care insurance market." *American Economic Review* 96 (4): 938–958.
- Folland, Gerald B. 1999. *Real Analysis: Modern Techniques and Their Applications*. Vol. 40. John Wiley & Sons.
- Friedman, Milton. 1962. *Price Theory*. Aldine Publishing Company.
- Fudenberg, Drew, and Jean Tirole. 1991. *Game Theory*. MIT press.
- Geruso, Michael, Timothy J Layton, Grace McCormack, and Mark Shepard. 2023. "The two-margin problem in insurance markets." *The Review of Economics and Statistics* 105 (2): 237–257.
- Hale, Jack K. 1980. *Ordinary Differential Equations*. Courier Corporation.
- Handel, Ben, Igal Hendel, and Michael D Whinston. 2015. "Equilibria in health exchanges: Adverse selection versus reclassification risk." *Econometrica: Journal of the Econometric Society* 83 (4): 1261–1313.
- Handel, Ben, and Kate Ho. 2021. "The industrial organization of health care markets." In *Handbook of Industrial Organization*, 5:521–614. 1. Elsevier.
- Handel, Benjamin R. 2013. "Adverse selection and inertia in health insurance markets: When nudging hurts." *American Economic Review* 103 (7): 2643–2682.
- Hazewinkel, Michiel, et al. 1990. "Encyclopaedia of Mathematics: Volume 5." *Springer Book Archive-Mathematics*.
- Hendren, Nathaniel. 2013. "Private information and insurance rejections." *Econometrica: Journal of the Econometric Society* 81 (5): 1713–1762.
- . 2014. "Unravelling vs unravelling: A memo on competitive equilibriums and trade in insurance markets." *The Geneva Risk and Insurance Review* 39:176–183.

- Hendren, Nathaniel. 2017. “Knowledge of future job loss and implications for unemployment insurance.” *American Economic Review* 107 (7): 1778–1823.
- Hicks, John R. 1937. “Mr. Keynes and the “classics”; a suggested interpretation.” *Econometrica: Journal of the Econometric Society* 5 (2): 147–159.
- Kallenberg, Olav. 2002. *Foundations of Modern Probability*. Springer Science & Business Media.
- Kamenica, Emir, and Matthew Gentzkow. 2011. “Bayesian persuasion.” *American Economic Review* 101 (6): 2590–2615.
- Kartik, Navin. 2009. “Strategic communication with lying costs.” *The Review of Economic Studies* 76 (4): 1359–1395.
- Keynes, John Maynard. 1936. *The General Theory of Employment, Interest and Money*.
- Leoni, Giovanni. 2017. *A First Course in Sobolev Spaces*. American Mathematical Society.
- Luce, R Duncan. 1959. *Individual Choice Behavior*. John Wiley.
- Maćkowiak, Bartosz, Filip Matějka, and Mirko Wiederholt. 2023. “Rational inattention: A review.” *Journal of Economic Literature* 61 (1): 226–273.
- Mailath, George J. 1987. “Incentive compatibility in signaling games with a continuum of types.” *Econometrica: Journal of the Econometric Society* 55 (6): 1349–1365.
- Mailath, George J, Masahiro Okuno-Fujiwara, and Andrew Postlewaite. 1993. “Belief-based refinements in signalling games.” *Journal of Economic Theory* 60 (2): 241–276.
- Mailath, George J, and Ernst-Ludwig Von Thadden. 2013. “Incentive compatibility and differentiability: New results and classic applications.” *Journal of Economic Theory* 148 (5): 1841–1861.

- Matějka, Filip, and Alisdair McKay. 2015. "Rational inattention to discrete choices: A new foundation for the multinomial logit model." *American Economic Review* 105 (1): 272–298.
- McFadden, Daniel. 2001. "Economic choices." *American Economic Review* 91 (3): 351–378.
- Milgrom, Paul, and John Roberts. 1982. "Limit pricing and entry under incomplete information: An equilibrium analysis." *Econometrica: Journal of the Econometric Society* 50 (2): 443–459.
- Mirrlees, James A. 1971. "An exploration in the theory of optimum income taxation." *The Review of Economic Studies* 38 (2): 175–208.
- Riley, John. 1979. "Informational equilibrium." *Econometrica: Journal of the Econometric Society* 47 (2): 331–59.
- Rothschild, Michael, and Joseph Stiglitz. 1976. "Equilibrium in competitive insurance markets: An essay on the economics of imperfect information." *The Quarterly Journal of Economics* 90 (4): 629–649.
- Shavell, Steven. 1979. "On moral hazard and insurance." *The Quarterly Journal of Economics* 93 (4): 541–562.
- Siegelman, Peter. 2003. "Adverse selection in insurance markets: An exaggerated threat." *The Yale Law Journal* 113:1223.
- Sims, Christopher A. 1998. "Stickiness." In *Carnegie-Rochester Conference Series on Public Policy*, 49:317–356. Elsevier.
- . 2003. "Implications of rational inattention." *Journal of Monetary Economics* 50 (3): 665–690.
- Spence, Michael. 1973. "Job market signaling." *The Quarterly Journal of Economics* 87 (3): 355–374.

- Teschl, Gerald. 2012. *Ordinary Differential Equations and Dynamical Systems*. Vol. 140. American Mathematical Society.
- Veiga, André, and E Glen Weyl. 2016. “Product design in selection markets.” *The Quarterly Journal of Economics* 131 (2): 1007–1056.