

Inducing Information Provision through Competition Policy: Prohibitions on False and Unsubstantiated Claims

Kenneth S. Corts*

University of Toronto

PRELIMINARY AND INCOMPLETE: PLEASE DO NOT CITE

January 6, 2009

Abstract

Competition law in many countries prohibits firms from making false claims about product quality or performance to potential buyers and also requires that the truth of specific claims be supported by adequate prior testing. This paper explores the differences between these two policies and asks, among other questions, whether a policy of mandatory prior substantiation has any incremental effect if a ban on false claims is in place. This paper develops a model in which firms have private information about their probability of having a high quality product and are able to determine product quality with certainty through costly learning. Penalties for false claims and for unsubstantiated claims create an opportunity for firms to credibly reveal their information and for signaling to emerge in equilibrium. I show that the two kinds of penalties affect the possibility of signaling in different ways, and that the mandatory substantiation requirement in many circumstances improves buyer information and social welfare beyond what is achieved by a ban on false claims alone. It is therefore not redundant to a false claims ban, but is a useful additional policy tool in markets characterized by asymmetric information.

1 Introduction

Competition law in many countries prohibits firms from making false claims about product quality and performance to potential buyers and also requires that specific performance claims be supported by adequate prior testing. This paper explores the differences between these two policies and asks, among other questions, whether a policy of mandatory prior substantiation has any incremental effect if a ban on false claims is in place. This is an especially timely question in Canadian competition policy; in a number of recent cases firms have mounted a defense to charges of unsubstantiated claims in part by arguing that the mandatory prior substantiation requirement should be held to be an unconstitutional infringement of free speech. This is based on an argument that this requirement is overly broad in the sense that any unsubstantiated claims that are false are already prohibited, and that unsubstantiated true claims are not problematic. Such a defense has succeeded in one recent case and failed in another. This is also a relevant policy question in the US, as a Federal

*Rotman School of Management, 105 St. George St., Toronto, Ontario M5S 3E6; kenneth.corts@rotman.utoronto.ca. I thank Joe Farrell and seminar participants at UC-Berkeley for helpful comments.

Trade Commission policy statement on the subject suggests ambivalence toward enforcement of the mandatory substantiation requirement, which is also present in US competition law.

This paper develops a model of firms with private information on product quality. Firms know their probability of having a high quality product and are able to determine product quality with certainty through costly learning. Penalties for false claims and for unsubstantiated claims create an opportunity for credible revelation of information, which may allow signaling to emerge in equilibrium. Throughout most of the analysis I assume that the combined effect of the penalties for false claims and unsubstantiated claims is sufficient to deter outright lying by firms: the penalties are assumed to be high enough that a firm that knows it definitely has a low quality product never finds it profitable to say it has a high quality product. The interesting contrast between the penalties comes about when instead an uncertain firm is considering making a “speculative claim”—that is, claiming it has a high quality product because it believes that this is likely, though not certain, to be true. Such a firm discounts the false claims penalty, sometimes significantly, because it believes there is a good chance the performance claim is true. In contrast, even a firm that is fairly certain it has a high quality product knows it is certainly violating the prior substantiation requirement by making a speculative claim, and thus does not discount that penalty. The mandatory prior substantiation penalty is therefore a more powerful tool for increasing the firm’s incentive to invest in learning product quality, especially when firms are optimistic about product quality. Speculative claims are appropriately discounted by consumers in this model, but since they contain less information than a substantiated claim of product quality, they do not achieve the full potential welfare gains associated with learning and full information revelation. Thus, the mandatory substantiation policy improves equilibrium information revelation and increases gross total welfare.

Sufficiently high penalties under either type of policy ensure equilibrium signaling in the model. However, using the false claims ban alone presents three problems, each of which is mitigated by the additional use of a mandatory prior substantiation policy. First, for a “small but reasonable” false claim penalty defined later (essentially, an expected penalty equal to the profits unduly gained through misrepresenting product quality), the false claims ban fails to achieve equilibrium signaling under any circumstances. The reason is that this penalty makes firms exactly indifferent between making the speculative claim, on one hand, and learning and telling the truth, on the other, but provides no incentive to bear a positive learning cost. Under a mandatory prior substantiation policy with a similar penalty, there is equilibrium signaling, the difference being that optimistic firms don’t discount this penalty; it therefore makes learning, from which they are fairly certain to receive good news, relatively more attractive for these firms.

Second, while higher false claims penalties expand the range of parameters for which equilibrium signaling arises, for any given false claims penalty there is a level of optimism such that the firm will prefer to make the speculative claim than learn its product quality. This occurs because a firm very sure it has a high quality product discounts the false claims penalty severely since it believes it is almost certainly making a true claim. Again, since these speculative claims contain less information than a fully informed product quality statement, this fails to maximize gross total welfare. A mandatory prior substantiation penalty is not discounted by such firms and encourages equilibrium signaling.

Third, higher penalties for false claims are successful in increasing the range of parameters for which there is equilibrium signaling; however, this includes ranges of parameters for which this is

socially beneficial and ranges for which it is socially wasteful. For the reasons described above, equilibrium signaling increases gross total welfare, but this must be weighed against the social cost of learning. Because optimistic firms discount the false claims penalty so dramatically, a very large penalty must be employed to support signaling; but, such penalties are seen as even more severe by a less optimistic firm. Thus, a less optimistic firm may be driven to socially wasteful learning by a penalty designed to encourage socially beneficial signaling by a more optimistic firm. Because the mandatory prior substantiation penalty generates a relatively more binding constraint for more optimistic firms, having this second lever available allows the policy-makers to target penalties in a way that strikes an appropriate balance between the social benefits of learning and its social costs.

Finally, the model illustrates an interesting trade-off that arises for either kind of penalty. In the model, consumers understand that when firms make speculative claims of high quality these must be discounted. This permits an equilibrium in which firms do not learn but make speculative claims, and consumers correctly infer that firms are signaling not product quality, but rather their belief of the chance that they have high product quality. This kind of signaling requires for internal consistency that firms do not learn their type. However, stronger penalties of either type increase the incentive to learn product quality, thereby reducing the willingness of firms to engage in this type of signaling, which can be the socially optimal form of signaling under certain circumstances.

This paper contributes to the literature by making the first explicit theoretical inquiry, to my knowledge, into the effects of mandatory substantiation policies. There is certainly a great deal of theoretical work on the effects of asymmetric information in markets (the literature on lemons problems and adverse selection), on strategies firms may undertake to credibly reveal their information (the literature on signaling and certification), and on policies governments may employ to improve information provision (the literature on mandatory disclosure and liability rules). However, there appears to be no existing treatment of mandatory substantiation policies, which are quite different from mandatory disclosure and other related and well-studied policies. While mandatory disclosure rules require that firms inform consumers of all information acquired through testing, mandatory substantiation requires that all information provided to customers be supported by testing. The former prohibits firms from hiding bad news, which generally reduces incentives to acquire information; the latter prohibits firms from making speculative claims, which generally increases incentives to acquire information.

Akerlof (1970) is the seminal citation in this literature. That paper lays out the basic model of informed sellers selling to uninformed buyers and establishes the iterative unraveling process by which consumers revise down their assessment of product quality as they consider which firms would be unwilling to sell at a price reflecting the average quality of the remaining firms. Viscusi (1978) and Grossman (1981) apply this same logic to a setting in which informed sellers can credibly reveal product quality; in this case the uninformed or pooled portion of the market (at least partially) unravels because high-quality firms reveal their information, lowering consumer expectations of average quality among pooled firms and inducing further information revelation. This leads to complete revelation of information if certification (i.e., a mechanism for credible information revelation) is costless, as in Grossman's model, or partial revelation if it is costly, as in Viscusi's model. Farrell (1986) and Matthews and Postlewaite (1985) pursue a similar analysis but focusing on firms that actively undertake learning to become informed. Both papers show that mandatory disclosure regulations dampen the incentive to acquire information and may obstruct the unraveling process and prevent complete

information revelation. A large literature on mandatory disclosure has since followed (including an offshoot literature in financial economics focusing specifically on financial disclosure regulations), including recent papers by Polinsky and Shavell (2006), who investigate how the liability regime affects the impact of mandatory disclosure on information revelation when firms must invest to learn product quality, and Daughety and Reinganum (2008), who analyze how the liability regime affects an exogenously informed firm’s decision of whether to signal its information through certification or through prices.

Section 2 provides context for this analysis by discussing the prevailing legal regime in Canada and the US. Section 3 lays out the basic model, which is analyzed in section 4. Section 5 puts more structure on the model to allow a somewhat more detailed welfare analysis. Section 6 concludes.

2 Background

In Canada, the existence of both a false claims prohibition and mandatory substantiation is quite clear. Section 74.01(1)(a) of the *Competition Act* prohibits any “representation to the public that is false or misleading in a material respect,” while section 74.01(1)(b) prohibits any “representation to the public . . . of the performance, efficacy, or length of life of the product that is not based on an adequate and proper test thereof. . .”¹ If the Competition Tribunal finds that a firm has violated these provisions, it may order that the firm cease making the suspect claims and/or engage in corrective advertising and it may levy administrative fines of up to \$100,000 for the first violation and \$200,000 per violation thereafter.

In the US, the situation is less clear, but the *FTC Policy Statement Regarding Advertising Substantiation* (which was issued in the 1980s but is prominently featured alongside other and more recent guidance documents) summarizes the FTC’s position.² This document states, seemingly quite clearly, that the FTC “intends to continue vigorous enforcement of this existing legal requirement that advertisers substantiate express and implied claims. . .” and that “... a firm’s failure to possess and rely upon a reasonable basis for objective claims constitutes an unfair and deceptive act or practice in violation of Section 5 of the FTC Act.” However, the document goes on to elaborate on the admissibility of post-claim evidence—that is, substantiating tests performed after the representations of quality or performance were made to the public—and states that “the truth or falsity of a claim is always relevant to the Commission’s deliberations. Therefore, it is important that the agency retain the discretion and flexibility to consider [post-claim] evidence. . .” The document proceeds to elaborate on the circumstances under which true but initially unsubstantiated claims will essentially be given the benefit of the doubt and not be pursued as violations of the FTC Act.

In Canada, the mandatory substantiation policy clearly exists as a separate requirement, but it is also under attack. In two recent cases brought by the government under section 74.01(1)(b), the firms have defended themselves in part through a constitutional challenge to the prior substantiation requirement. Such challenges are based on the free speech protections of the Charter. In broad terms, commercial speech is protected speech, and the provisions of section 74.01(1) are deemed to be infringements of that speech. However, the Charter permits the government to infringe on such rights in some circumstances if such infringement is necessary to advance certain valid policy

¹<http://laws.justice.gc.ca/en/ShowFullDoc/cs/C-34///en>

²<http://www.ftc.gov/bcp/guides/ad3subst.htm>

goals. To successfully justify the infringement, the government bears the burden of proof for certain tests pertaining to the value of the policy goals pursued and the nature of the infringement. Most relevant in these cases, the government must show that the infringement meets the test of “minimal impairment”—that is, that the same policy goal cannot readily be met with alternative restrictions that constitute a less significant infringement on the protected rights. The argument in these cases is essentially that section 74.01(1)(a) already prohibits false claims and that 74.01(1)(b) is therefore an overly broad prohibition—catching as it does true but unsubstantiated claims—that fails to meet the minimal impairment test.

A defense along these lines has been made in two recent cases heard by the Competition Tribunal. In *Commissioner of Competition v. Gestion Lebski* (2006 Comp. Trib. 32), the Tribunal found that the government had presented essentially no evidence to meet its burden of proof in justifying the infringement of free speech, and therefore held 74.01(1)(b) invalid for purposes of that case (but that case only). However, as the government had charged violations of both 74.01(1)(a) and 74.01(1)(b) in this case, Gestion Lebski was ordered to pay a fine for violation of section 74.01(1)(a). In *Commissioner of Competition v. Imperial Brush Co. and Kel Kem (c.o.b. as Imperial Manufacturing Group)* (2008 Comp. Trib. 02), the government did present evidence in support of its infringement of free speech, and the Tribunal held that 74.01(1)(b) was valid and applicable.³ Imperial Manufacturing Group was ordered to pay a fine for violation of 74.01(1)(b).

3 Model

This section lays out a simple one-shot game in which a firm with private information about product quality sells to consumers in the presence of a regulator that enforces penalties for certain prohibited behaviors.

3.1 A game of asymmetric information

A firm sells a product that may be of either high (H) or low (L) quality. A firm may be one of two types, where the type represents the firm’s probability of having a high quality product. The firm knows its type, which may be either 0, meaning that it is certain that it has a low-quality product, or γ , meaning that it has a $\gamma \in (0, 1)$ probability of having a high-quality product and a $1 - \gamma$ probability of having a low-quality product. The firm’s type is randomly drawn, with λ denoting the probability that a firm is type γ .

At the beginning of the game, the firm knows its type, but not its product quality. However, the firm can perform tests to learn its product quality at a cost of $c > 0$. Both the act of testing/learning and the resulting knowledge of product quality are unverifiable to consumers, but verifiable to the regulator upon investigation. The firm’s type is not verifiable to either consumers or the regulator, except inasmuch as it can be inferred with certainty after the fact in the case of a high-quality product.

Demand is higher for high-quality products than for low-quality products, and this results in higher profits for sellers of high-quality products. Throughout most of the paper, I employ reduced form expressions for profits, and do not specify the relationship of consumer surplus to product

³I provided expert testimony for the Commissioner of Competition in the constitutional challenge aspect of this case.

quality. Firms can make product quality claims to consumers, and consumers make inferences about true product quality based on these claims and their knowledge of the underlying distributions of types and qualities. The firm's profits when the consumers believe the firm's product is of quality x is π_x , where $x \in \{H, L\}$ denotes a belief that the product is of high or low quality respectively, and $x = \gamma$ denotes a belief that the firm of type γ and therefore that the product is of high quality with probability γ . I assume that $\pi_H > \pi_\gamma > \pi_L > 0$. I make one additional assumption that involves these profit expressions: $c \leq \pi_H - \pi_L$. This simply rules out scenarios in which the cost of learning is so large that the firm cannot under any circumstances justify paying this cost to improve its product quality in the eyes of consumers.

The final player in this game is the regulator, who plays a passive role and simply enforces penalties for prohibited behavior. There are two types of policies that prohibit different kinds of product quality claims made to consumers.

A “false claims ban” (FCB) prohibits statements of product quality that are in fact false. The exogenous and unspecified investigative and legal regime is such that firms that claim their product is high quality when in fact it is low quality face an expected penalty Δ_1 . Recall that product quality is verifiable by the regulator. Thus, the assumption here is that some exogenous fraction of claims are investigated and some exogenous penalty is levied when false statements are discovered. Since the penalty is levied after the firm has made its decisions about learning and made its product quality claims, only the expected penalty matters for the analysis.

A “mandatory prior substantiation” (MPS) policy prohibits claims that are not based on prior testing, regardless of the actual truth or falsity. Firms that claim high product quality but have not paid the cost c to learn their product's quality (and learned that it was in fact high quality) face an expected penalty Δ_2 . For the same reasons as before, only the expected penalty matters. A firm that paid to learn its product quality and learned its product was low quality, but which nonetheless claimed high quality, is subject to both penalties.⁴

3.2 Timing

The game proceeds in the following stages.

1. Nature chooses the firm's type γ , which the firm observes.
2. The firm can pay a cost c to learn its product quality (H or L).
3. The firm can make a claim of its product quality, stating H or L .
4. Consumers make inferences about product quality; demand and profits are realized.
5. The regulator assesses penalties for false and/or untested claims.

3.3 Signaling

The firm can attempt to signal its private information—be it about product quality or about its type—only through claims about its product quality. Absent the intervention of the regulator, there is no

⁴Both penalties could be thought of expansively to include legal costs and reputational costs incurred upon the regulator's finding of improper claims. These “penalties” could even include reputational or other costs incurred independent of the regulator's activities if, for example, there were other routes of consumer discovery of quality. For simplicity, I simply assume that the regulator controls this expected penalty.

differential cost of signaling that would allow signals to be credible. Since, moreover, the regulator can verify and therefore condition penalties on only the actual product quality and not firm type, only claims about product quality will generate the differential costliness (through differential exposure to penalties) that can ensure credibility of the signal. Thus, the restriction to signals based on product quality claims is without loss of generality. This is perhaps clearer in the case of a particular alternative signal.

The natural alternative to consider is that the firm signals its actual *ex ante* private information—that is, its type. Imagine that a firm does so, by truthfully claiming that “my product will be of high quality a proportion γ of the time.” One of two interpretations of the false claims ban must apply. One possibility is that this deemed to be an unverifiable claim and so not subject to the false claims ban (as construed in this model). In this case, there is no differential cost to the signal that would permit its credibility. Alternatively, this could be deemed to be a claim about actual product quality that is subject to penalties when false. In this case, the FCB penalty would apply the $1 - \gamma$ proportion of the time that the quality was not in fact high. In this case, this statement of type is equivalent to a statement of quality.

4 Analysis

This section develops the analysis of the main model. First, I lay out the basic inequalities, for arbitrary penalties, that determine when a signaling equilibrium exists. I consider two kinds of signaling equilibria—one in which consumers infer firm type, and one in which consumers infer product quality. I then present analytical results for the extreme cases of a false claims ban only and a mandatory prior substantiation policy only. Finally, I present a graphical analysis that permits both policies to be in force simultaneously.

4.1 Weak and strong separation

I examine two kinds of signaling equilibria. One is what I call a “weak separation” equilibrium. This is a standard signaling equilibrium in which consumers infer firm type from signals. This is “weak” because, while it does result in symmetric information, it does not result in perfect information, as both firms and consumers remain unsure of product quality. The other is what I call a “strong separation” equilibrium. This is a standard signaling equilibrium in which consumers infer product quality from signals. This is “strong” in that it results in both symmetric and perfect information; it is the situation in which consumers actually know what product quality they are buying.

Both of these are standard signaling equilibria in that I require incentive compatibility for firms and consistency of inferences by consumers. In principle, three kinds of incentive compatibility constraints apply: participation constraints, constraints that ensure the appropriate amount of firm learning for that equilibrium, and incentive constraints that ensure truthful signaling conditional on the extent of learning. In fact, the participation constraints are moot in this analysis since by assumption the firm makes positive profit even under the most pessimistic consumer inferences.

I first lay out the inequalities that are required to support the weak separation equilibrium. Recall that in this equilibrium, consumers infer that a claim of high quality implies the firm is type γ and that the absence of such a claim implies the firm is type 0, and that firms make high quality statements if and only if they are type γ . Thus, the firm in this case is making what I will call a

“speculative claim”—that is, they are claiming high product quality when in fact they know only that this is a possibility, not a certainty.

The incentive compatibility constraints that ensure truth-telling conditional on knowing only one’s type are:

$$IC_W(\gamma) : \pi_\gamma - \pi_L \geq (1 - \gamma)\Delta_1 + \Delta_2$$

$$IC_W(0) : \pi_\gamma - \pi_L \leq \Delta_1 + \Delta_2.$$

The left-hand side, $\pi_\gamma - \pi_L$, is the gross benefit to signaling. The right-hand sides of these inequalities represent the cost of signaling. In both cases the firm is exposed to the expected MPS penalty Δ_2 since it has not learned. The FCB penalty Δ_1 is added to the cost of signaling if the firm is type 0, but for the type γ firm this is discounted by the probability $1 - \gamma$ that the claim is false. This pair of incentive constraints simply ensures that signaling pays for the high type firms but not for the low type firms; this can hold because, though the benefit to signaling is the same for the two types, the exposure to penalties for false claims differs between the two types of firms.

For these to be the right incentive constraints, it must be that the firm chooses not to learn its type. This leads to a “learning constraint” that ensures the firm does not have an incentive to learn, which would allow it to avoid making speculative claims, thereby eliminating its exposure to the FCB and MPS penalties:

$$LC_W(\gamma) : \pi_\gamma - (1 - \gamma)\Delta_1 - \Delta_2 \geq \gamma\pi_\gamma + (1 - \gamma)\pi_L - c.$$

The left-hand side here represents the net profit of making the speculative claim. The firm is identified as a type γ but is subject to both types of penalties. The right-hand side reflects the alternative net profit the firm would earn by learning. It is clear that a firm that learned it had a low-quality product would not signal, since this is precisely the condition given in $IC_W(0)$. Thus, the profit that follows from learning is the weighted average of the profits that accrue to a firm who has revealed itself as a type γ or a type 0 firm, less the cost of learning. Such a firm has no exposure to either of the penalties.

A similar set of inequalities determines the set of parameters for which strong separation can be sustained. Recall that in this equilibrium, consumers infer that a claim of high quality implies the firm has a high quality product and that the absence of such a claim implies the firm has a low quality product, and that the firm makes the high quality claim if and only if it is in fact high quality. For this to be true, of course, the firm must pay to learn its product quality. This gives rise to a similar set of incentive and learning constraints.

The incentive constraints that ensure truth-telling conditional on knowing product quality are:

$$IC_S(H) : \pi_H - \pi_L \geq 0$$

$$IC_S(L) : \pi_H - \pi_L \leq \Delta_1 + \Delta_2.$$

As in the case of the weak separation incentive constraints, the left-hand side represents the gross benefit to signaling, while the right-hand side represents the cost of expected penalties. Note that a

firm that has learned it has a high-quality product faces no exposure to either kind of penalty, while the firm who has learned that it has a low-quality product faces both expected penalties.

For these to be the right incentive constraints, it must be that all type γ firms pay to learn their true product quality. This generates in this case two distinct learning constraints, one that requires learning to be more profitable than a speculative claim, and a second that requires learning to be more profitable than simply not signaling at all:

$$LC_S^1(\gamma) : \gamma\pi_H + (1 - \gamma)\pi_L - c \geq \pi_H - (1 - \gamma)\Delta_1 - \Delta_2$$

$$LC_S^2(\gamma) : \gamma\pi_H + (1 - \gamma)\pi_L - c \geq \pi_L.$$

In both cases, the left-hand side represents the profit earned by learning and then truthfully signaling. In LC_S^1 , the right-hand side represents the profits from making a speculative claim, including the exposure to penalties. In LC_S^2 , the right-hand side reflects the profits the firm earns by simply not signaling and being viewed as a seller of a low-quality product.

4.2 Separating and pooling equilibrium existence

These sets of inequalities determine the set of parameters for which each type of signaling equilibrium exists. That is not to say that such an equilibrium is the only equilibrium for a given set of parameters that satisfies those constraints. Rather, given the fact that the participation constraints are trivially satisfied, it is straightforward that for any set of parameter values there exists a pooling equilibrium in which no firm signals; in such an equilibrium, consumers infer nothing from statements about product quality but believe that the firm's product is high quality with probability $\lambda\gamma$. That said, it is true that there is at most one type of signaling equilibrium for any set of parameter values. This follows immediately from the fact that LC_S^1 and LC_W cannot be simultaneously satisfied. Rearranging and combining these two constraints yields $\pi_H - \pi_L \leq \Delta_1 + \frac{\Delta_2 - c}{1 - \gamma} \leq \pi_\gamma - \pi_L$, which contradicts the assumption that $\pi_H > \pi_\gamma$.

Proposition 1 *For any set of parameter values, there exists a pooling equilibrium and at most one type of separating equilibrium.*

One other general observation that can be made about the existence of a separating equilibrium is that no signaling equilibrium of either type will exist in the absence of regulatory penalties. It is only these penalties that create the differential cost of signaling that permit the signal to be credible and prevent imitation by low type firms or firms with low-quality products. This follows immediately from the fact that $IC_W(0)$ and $IC_S(L)$ cannot hold at $\Delta_1 = \Delta_2 = 0$, given the assumption that $\pi_H - \pi_L > \pi_\gamma - \pi_L > 0$.

Proposition 2 *For any set of parameter values, no separating equilibrium of either type exists in the absence of regulatory intervention.*

4.3 False claims ban only

In this subsection I derive analytical results for the case in which there is a false claims ban but no mandatory prior substantiation. This illustrates the strengths and weaknesses of this kind of

regulation in achieving separation. It also establishes the baseline case against which one can ask the policy question of whether mandatory substantiation provision has incremental value.

I emphasize throughout the paper that the false claims ban has three primary weaknesses. Two of these can be explicitly addressed within the model and form the basis of the next two propositions: FCB alone supports no strong separation for small but reasonable penalties; and, no matter how high the FCB penalty or how low the learning cost, FCB alone supports no strong separation for a firm very confident in its product's high quality. The third weakness—that increasing FCB penalties to achieve strong separation encourages socially wasteful excess learning—requires more structure in order to address the net social value of separation and must wait for a later section.

Proposition 3 *For “small but reasonable” penalties ($\Delta_1 = \pi_H - \pi_L$), a false claims ban alone achieves no strong separation, regardless of the other parameter values.*

Note that this “small but reasonable” penalty is exactly the “unfair profit” gained by a firm that falsely claims to be high quality under the consumer inferences associated with strong separation. By “reasonable” I do not mean to imply that this is a penalty that constitutes a good policy by some measure. I mean only that it is not a pathological case of arbitrarily high or low penalties, but is a penalty that is related in a reasonable way to the model's parameters. The penalty is reasonable in the sense that it does not entail large punitive damages but essentially taxes away the profit gains achieved through the false quality claims. To the extent that the firm's profit increase represents a pure transfer of consumer surplus, this penalty is also a good proxy for the harm suffered by consumers, which could be relevant under certain liability rules that might constrain the size of the penalty. The penalty is small in the sense that it is the minimal penalty that satisfies $IC_S(L)$, the informed, low-quality firm's incentive constraint. It therefore permits strong separation with respect to that constraint. The content of the proposition lies in showing that a penalty that just satisfies this incentive constraint strictly fails to satisfy the strong separation learning constraint LC_S^1 .

That constraint can be rearranged in the FCB-only case ($\Delta_2 = 0$) to yield $\pi_H - (1 - \gamma)(\pi_H - \pi_L) - c \geq \pi_H - (1 - \gamma)\Delta_1$. This can be read as follows: the downside to learning (on the left-hand side) is that $1 - \gamma$ of the time the firm gets only π_L and the firm spends c , while the downside of making the speculative claim (on the right-hand side) instead is that $1 - \gamma$ of the time the firm is exposed to the FCB penalty. As this inequality makes clear, at these small but reasonable penalties ($\Delta_1 = \pi_H - \pi_L$) the firm is indifferent between the risk of discovering it is a low type and the risk of paying the FCB penalty. However, the firm strictly prefers not to learn since learning is costly. With the failure of the LC_S^1 constraint, there can be no strong separation, which proves the proposition.

The next proposition highlights the second weakness of a false claims ban alone—its ineffectiveness in eliminating speculative claims by firms very confident in their high product quality even when penalties are very large.

Proposition 4 *No matter how high the penalty or how low the learning cost, a false claims ban supports no strong separation if the firm is very confident that it has a high quality product. (For any Δ_1 and any c , there exists some $\hat{\gamma} < 1$ such that no strong separation equilibrium exists for $\gamma > \hat{\gamma}$.)*

The intuition for this result is very simple. The incentive to learn in this model derives from the opportunity to avoid making a false quality claim that exposes the firm to FCB penalties. However, a firm that is very confident about its product quality prior to learning (that is, has a very high γ)

discounts the possibility of being subject to the penalty (in the limit, it ignores the penalty). This leaves the firm no incentive to learn, which invalidates the consumers' inference under strong signaling that high quality claims are made only by firms who have high quality with certainty. Algebraically, there is no strong signaling for any parameters if $\Delta_1 < \pi_H - \pi_L$ (by the incentive constraint $IC_S(L)$); for larger Δ_1 , the learning constraint LC_S^1 is the relevant constraint. This can be rearranged to yield $c \leq (1-\gamma)[\Delta_1 - (\pi_H - \pi_L)]$, both sides of which are continuous in γ . In the limit as $\gamma \rightarrow 1$, this becomes $c \leq 0$, which is clearly strictly false. It is therefore strictly false for γ near enough 1. (Technically, the learning constraint is strictly false for $\gamma > \hat{\gamma} = \max\{0, 1 - c/[\Delta_1 - (\pi_H - \pi_L)]\}$.) Without satisfaction of the learning constraint, there is no strong separation, which proves the proposition.

4.4 Mandatory prior substantiation only

Mandatory prior substantiation and the false claims ban each have their strengths and weaknesses, which will be compared in the next subsection. Since the main emphasis of the paper is that the two policies can complement each other, I do not want to devote too much of the analysis to the individual policies alone. However, it is useful for analytical clarity to show in this subsection how mandatory prior substantiation on its own is less prone to the problems demonstrated above for a false claims ban alone. The next two propositions provide the counterpoint to the prior two.

Proposition 5 *For “small but reasonable” penalties ($\Delta_2 = \pi_H - \pi_L$), mandatory prior substantiation alone achieves strong separation if learning costs are low enough (specifically, if $c \leq \gamma(\pi_H - \pi_L)$), regardless of the other parameters.*

These are the same small but reasonable penalties as in the earlier proposition, which showed that FCB alone failed to achieve strong separation. Again, this penalty is the smallest that satisfies the informed, low-quality firm's incentive constraint, and it therefore permits strong separation with respect to that constraint. Recall that the earlier proposition showed that it therefore failed to satisfy the learning constraint for the high type (LC_S^1) and therefore did not support strong separation. The content of this proposition is in showing that this same penalty, when applied to unsubstantiated claims instead of false claims, does satisfy the learning constraint and permit strong separation. Employing the analogous rearrangement of LC_S^1 in this case yields $\pi_H - (1-\gamma)(\pi_H - \pi_L) - c \geq \pi_H - \Delta_2$. In the left-hand side, which represents the profits from learning and truthfully signaling, the firm accounts for the $1 - \gamma$ chance that it will discover it does not have a high-quality product, as well as the learning cost it will pay. In the right-hand side, the firm accounts for the substantiation penalty it will pay if it does not learn but signals anyway. When $\Delta_2 = \pi_H - \pi_L$, this inequality is strictly satisfied for c near 0 (specifically for $c \leq \gamma(\pi_H - \pi_L)$) since the $\pi_H - \pi_L$ on the left-hand side is discounted by $1 - \gamma$. The penalty for unsubstantiated claims, which is certain, looms very large in comparison to the chance that the firm will learn it has a low quality product and choose not to signal.

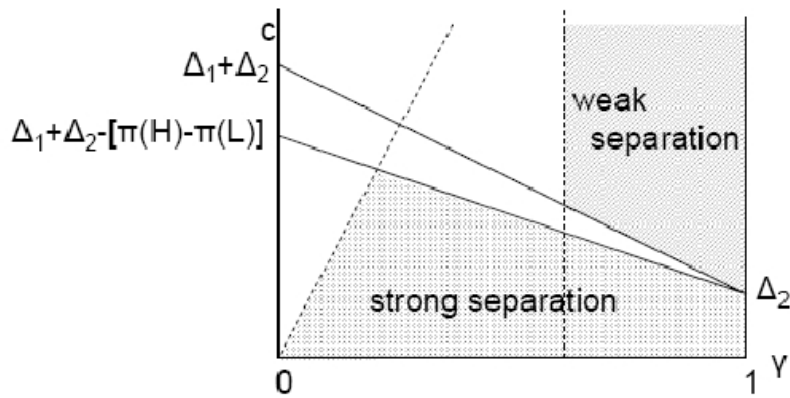
Proposition 6 *For any penalty at least as large as the small but reasonable penalties considered above, there exists a range of learning costs $c < \Delta_2$ for which mandatory substantiation alone supports strong separation even if the firm is very confident in its product's high quality (specifically, for any $\gamma \in (\hat{\gamma}, 1)$, where $\hat{\gamma} = \max\{\frac{c}{\pi_H - \pi_L}, 1 - \frac{\Delta_2 - c}{\pi_H - \pi_L}\}$).*

The analogous proposition in the prior subsection showed that false claims penalties alone could not support strong separation for firms that were very confident in the product's high quality. This proposition shows that mandatory substantiation can. The intuition is simple. The FCB penalty fails to support strong separation because the penalty is discounted by the small probability that the firm making a speculative claim will be exposed to that penalty. In the case of the unsubstantiated claim penalty, there is no such discounting. By making a speculative claim, the firm is exposing itself to the MPS penalty regardless of its level of confidence in its product quality. In the limiting case, the firm is comparing learning, followed by truthfully signaling its high quality (incurring only the learning costs), with signaling without learning (exposing itself to the MPS penalty). If the learning cost is less than the MPS penalty, then learning is preferable. For learning costs strictly lower than the MPS penalty, a firm somewhat less confident of its high quality will similarly prefer to learn rather than to expose itself to the MPS penalty. The threshold γ given in the proposition is determined by the binding learning constraint for that particular learning cost, either LC_S^1 or LC_S^2 .

4.5 Combined policies

While the propositions of the last two subsections deal with the separation regimes that prevail when only one type of policy is in place, the inequalities derived in subsection 4.1 fully determine the separating equilibria that may arise when a combination of the two policies is used instead. The most effective way to understand the effects of the combined policies, and the differential impact of increasing the two penalties, is to plot the relevant constraints and examine in a figure the regions in which separation of either type can be sustained.

The following figure depicts these constraints and these regions in γ, c space, holding the profit functions fixed. Since no strong separation can occur unless the low type's incentive constraint $IC_S(0)$ is satisfied (under the consumer expectations associated with strong signaling), I assume in drawing this figure that the penalties of the policies are together sufficiently large to enable to support strong signaling for some c, γ : $\Delta_1 + \Delta_2 \geq \pi_H - \pi_L$.



The assumption on the minimal penalties ensures that $IC_S(L)$ holds, while $IC_S(H)$ holds by the original assumptions on π_i . In this figure, the line through the origin represents the $LC_S^2(\gamma)$ constraint: below it, learning and signaling (if the firm finds its product is high quality) is preferred to simply not signaling. Note that the payoffs involved in this trade-off, and therefore this line in the

figure, do not in any way depend on the size of the penalties. The lower of the dashed lines (note that it need not be down-sloping) represents the $LC_S^1(\gamma)$ constraint: below it, learning and signaling (if the firm finds its product is high quality) is preferred to not learning but making a speculative high-quality claim. Thus, below these two lines, in the shaded region labeled “strong separation,” all the relevant constraints hold and strong separation is supported by the given penalties.

The only constraint depicted here for the strong separation case that varies with the magnitude of the penalties is the $LC_S^1(\gamma)$ constraint. Note the intercepts of this constraint at $\gamma = 0$ and $\gamma = 1$. The left intercept is determined by the sum of the penalties, while the right intercept is determined by only Δ_2 . The intuition for this is simple. Toward the left side of the figure—that is, with respect to firms that know they are almost certain to have a low quality product—the two penalties are nearly perfect substitutes. Making a speculative claim is very likely (in the limit, certain) to lead to paying the FCB penalty Δ_1 , while it is certain to lead to exposure to the MPS penalty Δ_2 . In contrast, at the right side of the figure—that is, for firms that know they are almost certain to have a high quality product—the FCB penalty Δ_1 is of little concern.

It is possible to visualize propositions 3-6 using this figure. Proposition 3 can be visualized in this graph by noting that when $\Delta_1 + \Delta_2 = \pi_H - \pi_L$ the $LC_S^1(\gamma)$ constraint defining the upper boundary on strong separation is coincident with the horizontal axis under FCB only (i.e., when $\Delta_2 = 0$ and therefore $\Delta_1 = \pi_H - \pi_L$). Thus, no strong separation can occur for positive c . Similarly, under FCB only (i.e., when $\Delta_1 > 0$, $\Delta_2 = 0$), this line intersects the horizontal axis at $\gamma = 1$. The area above this line and near $\gamma = 1$ is the problematic region, described in Proposition 4, in which FCB-only fails to achieve strong separation even when $\Delta_1 + \Delta_2 > \pi_H - \pi_L$, which yields a positive intercept at the left axis.

Similarly, the MPS-only results can be visualized in this figure. Proposition 5 provides the counterpoint to proposition 3: in graphical terms, it states that even when this $LC_S^1(\gamma)$ constraint goes through the origin there will still be a region of strong separation since even a high- γ firm cares about the MPS penalty and this line is thus up-sloping. Proposition 6, which is analogous to proposition 4, follows from the fact that this constraint has a positive $\gamma = 1$ intercept under MPS, and thus there is necessarily a region of strong separation near the lower right corner of this graph that extends arbitrarily close to the $\gamma = 1$ axis.

More importantly, one can use this figure to think about the effect of combining these policies with different penalties. Again the emphasis is on the $LC_S^1(\gamma)$ constraint, which defines the upper bound of the strong separation region. The simplest way to think of this is that the height of this constraint can be fixed through the MPS penalty Δ_2 , while the slope is determined by the FCB penalty Δ_1 . The line is flat when the FCB penalty alone just satisfies the low type’s incentive constraint $IC_2(L)$ —i.e., when it takes the “small but reasonable” value employed in the propositions. Raising the FCB penalty above that level makes this line down-sloping, while lowering it below this level makes it up-sloping. If one fixes the sum of the penalties but shifts the penalty toward MPS, the line rotates up (counterclockwise) around its $\gamma = 0$ intercept.

This figure also depicts the region in which these policies support weak signaling. Since $\pi_H > \pi_\gamma$, the $IC_W(0)$ constraint is satisfied by the assumption employed in the figure that $\Delta_1 + \Delta_2 \geq \pi_H - \pi_L$. The vertical line that determines the lower γ boundary for weak signaling is the $IC_W(\gamma)$ constraint. The intuition here is straightforward: firms must be sufficiently hopeful that they will have a high-quality product in order to be willing to risk the FCB penalty that accompanies the speculative claim

that characterizes weak signaling. Increases in the penalties shift this line to the right. Moreover, given a fixed sum of the penalties, shifts toward MPS penalties shift the line right. The lower c boundary of this region is the $LC_W(\gamma)$ constraint. It is *not* a straight line as depicted in this graph. It does have the endpoints depicted here, but for $\gamma \in (0, 1)$ the constraint itself lies strictly between the line depicted and the lower dashed line representing the strong signaling learning constraint. This follows from the fact that $\pi_L < \pi_\gamma < \pi_H$. This line does lie everywhere above the strong signaling learning constraint, and it responds qualitatively similarly to changes in the composition of Δ_1 and Δ_2 . Specifically, increases in Δ_1 make the line more down-sloping, while increases in Δ_2 shift the line upwards. Fixing the sum of the penalties, putting more emphasis on the MPS policy tends to rotate this constraint around its $\gamma = 0$ intercept and make it more up-sloping.

Inspection of this figure, given this understanding of how the constraints change with the penalties, leads to several general conclusions about the differential effects of the two policies on the kinds of equilibria that can be supported. Determining how these should be weighed against each other depends on the social costs and benefits of these changes, a topic that is addressed in the next section. For now, I simply gather some of the observations (which can also be proved algebraically) from this discussion in the next proposition.

Proposition 7 *Given some penalties such that a firm would not knowingly misrepresent its type ($\Delta_1 + \Delta_2 \geq \pi_H - \pi_L$), increases in the substantiation penalty Δ_2 —either holding Δ_1 fixed or holding $\Delta_1 + \Delta_2$ fixed—increase the range of parameters for which strong separation is possible and decrease the range of parameters for which weak separation is possible. If $\Delta_1 + \Delta_2$ is fixed, increases in Δ_2 increase the range of learning costs for which strong separation is possible more quickly for firms that are more confident they have a high quality product.*

5 Optimal Policies

It is clear from the previous section that these two policies have different effects on the ability to achieve strong and weak separation. Ideally, one would like to know how to use these policies together to accomplish a particular policy goal such as, perhaps, maximizing total social welfare. Given the development of the model thus far, this is impossible since the generality of the model precludes calculation of the social benefit to learning and signaling. In this section I discuss some of the social benefits in general terms and then parameterize one particular model that will allow a suggestive analysis of the optimal combination of policies.

5.1 The social costs and benefits of weak and strong separation

There are two broad types of arguments for a social value to separation: that separation is crucial to quality or innovation incentives, and that separation generates a short-run total welfare increase because of convexity of the total social welfare function in quality. To provide context for this discussion, I first lay out a simple standard model in which separation has no social benefit.

Assume that a single price-setting seller has a single unit of a good to sell, that the firm has no outside option (i.e., a 0 reservation price or opportunity cost), and that there are many identical risk-neutral consumers, each of whom value the good at 10 if it is high quality and 6 if it is low quality. Assume further that the firm has no influence over product quality, that the firm knows its

product quality with certainty, and that the true ex ante probability of the firm having a high quality product is 50%. Without separation, the consumers' expected valuation is 8. The firm can charge a price of 8 and sell its one unit. In this case, consumer surplus is 0, producer surplus is 8, and total welfare is 8. Assume that separation is costlessly possible. Then if the firm has a high quality product it sells the unit at 10; if not, it sells the unit at 6. Expected consumer surplus is 0, expected producer surplus is $.5(10)+.5(6)=8$, and total welfare is $.5(10)+.5(6)=8$. In this case, even though separation is costless, it is of no social value. There is a distributional effect of separation in that, for example, a consumer fares worse and the firm fares better under separation than under pooling if the product is actually high quality. However, there is no effect on (expected) total social welfare. There are at least three crucial assumptions in this model that drive this result.

The first assumption that can be altered to yield a social value to separation is that product quality is exogenous. Suppose instead that the firm must pay a small cost in order to become high quality, reinterpreting the 50% as the probability that the firm has this opportunity to expend a small cost to improve its product quality. Without separation, the firm cannot expect to improve its profit through improving its quality since it will achieve the pooling profit regardless of its quality. It therefore has no incentive to invest in improving product quality. Endogenous product quality, which is surely a feature of virtually any real world market, therefore establishes a social value to separation even absent convexity of the social welfare function in quality.

There are at least two ways to plausibly alter assumptions to create a short-run social value to separation by introducing convexity in the social welfare function. The first of these is to assume that not all qualities of the product are "in the money"—that is, not all qualities have willingness-to-pay greater than opportunity cost, or not all qualities are "good trades." In the example above, this would mean assuming a reservation price/opportunity cost of strictly greater than 6. If, moreover, the reservation price/opportunity cost lies below the expected willingness to pay under pooling—in this case, less than 8—then pooling leads to a different and inferior set of trades being made than does separation. Pooling leads to both trades being made since the expected willingness to pay exceeds the cost, despite the fact that trade of the low quality good destroys value. Under separation, there is no trade of the low quality good, which makes the social welfare function convex in quality. This in turn makes the expected value of the two separation outcomes preferable to the value of the pooling equilibrium (which is evaluated at the expected value of quality). Note that in this example consumer surplus is zero in all cases, and producer surplus always equals total surplus. To put specific numbers on the effect described here, assume that the firm's opportunity cost of the good is 7. Then, under pooling, the trade is made regardless of realized quality, and $PS=TS=8-7=1$; under separation, the trade is made only under a high quality realization, and $PS=TS=.5(10-7)+.5(0)=1.5$.

The second alteration to the model that will generate total surplus that is convex in quality, and therefore a social gain to separation over pooling, is to change the unit demand assumption. If the firm instead faces down-sloping demand but must set a single price (i.e., the classical monopoly model), then it will choose a higher price and a higher output level when quality is higher. If the gain in surplus relative to pooling when quality is high is larger than the loss in surplus relative to pooling when quality is low, then surplus is convex in quality and there is a social benefit to separation. This is in fact true in the canonical linear-demand monopoly model, as detailed in the next subsection.

Finally, note that these social gains must be weighed against the social costs of separation. The social cost of separation in the model of this paper is simply the cost of learning. Evaluating this

trade-off requires a model that allows a specific characterization of the social gain to separation and how this varies with the parameters of the model. By focusing on a down-sloping demand model that generates a social benefit to separation, the next subsection provides the foundation for a subsequent analysis of when separation is socially optimal.

5.2 A model of the social value of separation

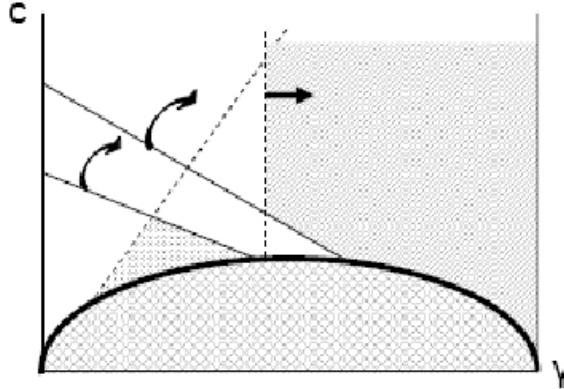
I now introduce a specific demand model to the basic model of sections 2 and 3 in order to provide structure on the π_i functions that figured prominently in that analysis. Specifically, suppose that the firm has constant marginal cost of zero and sets a uniform price to risk-neutral consumers whose demand curve is given by $P = a_x - Q$, where $x = H, L$ represents the consumer's belief of product quality. Let $a_x = a + \delta I_H$, where δ is the increment in willingness-to-pay for high quality products over low quality products and I_H is an indicator variable that is equal to 1 if the product is high quality and 0 otherwise. Since the consumers are risk-neutral, the demand curve when consumers believe there is a γ probability that the firm's product is high quality is $P = a + \gamma\delta - Q$; thus, in a slight abuse of notation, I will let a_γ denote $a + \gamma\delta$. This will be the relevant perceived quality under weak signaling.

The standard profit-maximization problem yields an optimal price and quantity of $p_x = q_x = \frac{a_x}{2}$. It is straightforward to calculate the relevant welfare measures. Gross producer surplus (not accounting for any learning cost) is $PS_x = \frac{a_x^2}{4}$, and consumer surplus is $CS_x = \frac{a_x^2}{8}$. I will use the term gross total surplus (denoted TS) to refer to the sum of the gross producer and consumer surplus, and the term net total welfare (denoted TW) to refer to this gross surplus less the learning cost, if any. I will use the total surplus and total welfare measures to compare different types of equilibria, and I will therefore subscript them with P , W , or S , to denote pooling, weak separation, and strong separation; such comparisons take into account the proportion of high and low types (i.e., λ). This analysis already implies the convexity of the welfare measures in the quality of the product, as the quality is a linear component of a_x , which enters the welfare measures as a quadratic. This is what creates the social value of information in this model.

The point of this exercise is that quantify the social benefit to weak and strong separation, allowing comparison of these benefits to the cost. It is straightforward to calculate, given the analysis above, expressions for the ex ante gross total surplus under each of the three regimes, given λ and γ : $TS_P = \frac{3}{8}a_{\lambda\gamma}^2$; $TS_W = \frac{3}{8}[\lambda a_\gamma^2 + (1 - \lambda)a_L^2]$; $TS_S = \frac{3}{8}[\lambda\gamma a_H^2 + (1 - \lambda\gamma)a_L^2]$. Some algebra shows that weak separation has higher gross total surplus than pooling: $TS_W - TS_P = \frac{3}{8}\lambda(1 - \lambda)(\gamma\delta)^2 > 0$. Since weak signaling involves no expenditure of the learning cost, note that this also implies an identical comparison of net total welfare under the two regimes, and weak separation is always socially preferred to pooling. Similarly, one can show that strong separation has higher gross total surplus than weak separation: $TS_S - TS_W = \frac{3}{8}\gamma(1 - \gamma)\lambda\delta^2 > 0$. Note that this does not necessarily imply that net total welfare after learning costs is higher under strong separation. In net total welfare terms, strong separation is preferred only if $TW_S - TW_W = \frac{3}{8}\gamma(1 - \gamma)\lambda\delta^2 - \lambda c = \lambda[\frac{3}{8}\gamma(1 - \gamma)\delta^2 - c] > 0$. Note that this always holds for large enough δ —that is, for a large enough impact of the high quality product on willingness to pay. Also, it is always negative for γ close enough to 0 or 1. In these cases the firm's type is essentially sufficient to determine the product quality, so weak separation achieves essentially all of the social benefits of separation, but without expenditure of the learning cost, rendering it preferable.

The structure of this model allows a characterization of the total welfare maximizing policy for any set of parameters by evaluation of the above inequalities. As described above, pooling is never preferred; either weak separation or strong separation is the socially optimal regime, depending on the parameters. Given a set of parameters, implementing the welfare-maximizing equilibrium is simple. The weak separation equilibrium can be implemented for any c, γ by, for example, using only the false claims ban, with a penalty of $\Delta_1 = \pi_\gamma - \pi_L + \epsilon$, for small ϵ . If instead the strong separation equilibrium is preferred, this can be implemented by, for example, using only the mandatory prior substantiation policy with a large enough penalty Δ_2 , as long as $\pi_H - \pi_L > \frac{c}{\gamma}$ (otherwise, the learning cost is too high for strong separation to ever be possible).

What is complex about determining the optimal policies is not how to implement them to achieve the welfare-maximizing equilibrium for a specific set of parameters, but rather how to design the policies to support the welfare-maximizing equilibrium over a range of parameter values. In the context of the earlier figures graphed in γ, c space, the $TW_S - TS_W > 0$ inequality derived above can be graphed to determine the regions of parameters for which weak and strong separation are preferred. This illustrates the difficulties associated with achieving the total welfare maximizing outcome. The following figure portrays the region in which strong separation maximizes net total welfare.



The parabola represents the $TS_S - TS_W > 0$ inequality derived above; below it, net total welfare is maximized through strong separation, while elsewhere weak separation is preferable. This is overlaid on the FCB-only graph; the arrows represent the effects of increasing the FCB penalty Δ_1 . This illustrates the third weakness of FCB alone discussed earlier. Ratcheting up the penalty to enlarge the region of strong separation also enlarges the area of inefficient learning—that is, socially wasteful learning for parameters under which weak separation is preferred to strong separation. In the figure, this is the shaded area just outside the parabola on the left side of the graph. This also illustrates one of the strengths of MPS as an alternative policy. Because it hits all types of firms equally hard for speculative and untested claims, the region of strong separation is more uniform across γ . Recall that in the figure a transfer of penalty from Δ_1 to Δ_2 rotates the strong separation constraint counterclockwise, expanding the region of strong separation under the parabola (where it is socially inefficient) more quickly relative to the increase in the area of inefficient signaling than does an increase in the FCB penalty Δ_1 .

Fully characterizing the optimal policy for some limited circumstances within the context of this

model remains work in progress. What can be readily conjectured, though, is that the optimal policy will typically entail some use of both false claims penalties and mandatory prior substantiation penalties since this will allow the policies to support strong separation over some desirable range of parameters without inducing socially wasteful learning over excessively large other ranges of parameters.

6 Conclusion

Prohibitions on false claims and prohibitions on unsubstantiated claims can have quite different effects in markets where firms have private information but are uncertain of their product quality. In such settings, inducing the firms to learn more about their product quality and then to truthfully reveal this to consumers is a total welfare enhancing policy goal if learning costs are not too large and/or there is wide variation in product quality. If, instead, firms make speculative claims about product quality, consumers still purchase under imperfect information, which leads to total welfare losses in many circumstances.

A ban on false claims alone suffers three shortcomings. First, it supports no learning if penalties are limited to the amount of unfair profits earned through false claims. Second, even with the penalty increased beyond this level, there are always levels of optimism such that a firm will not learn because it discounts the penalty for false claims so completely. Third, increases in the penalty for false claims to a level sufficient to induce optimistic firms to learn leads to socially wasteful expenditures on learning by pessimistic firms. Using mandatory substantiation in conjunction with a false claims ban provides more flexibility and allows the policy-maker to mitigate the weaknesses of the false claims ban.

It is worth discussing here how one should interpret in practical applications the model's distinction between product quality and firm type. In the model, the key distinction is that product quality is an objective fact about the product that is testable and verifiable to the policy enforcer, while the firm's type is a subjective assessment, which is not testable and not verifiable to the policy enforcer, that the firm has about its likely product quality. It is this distinction in testability and verifiability that is critical, not the distinction between certainty (H or L) and uncertainty (a probability γ in a range from 0 to 1) that is also present in the model. Thus, a product quality claim could absolutely involve uncertainty—for example, a claim that an herbicide for lawns will kill 80% of weeds on average. While this claim involves uncertainty, it is a testable claim, and it is a claim that could be verified by or to the policy enforcer upon investigation. It could therefore be a high quality claim in the context of the model. A firm's type in this context would be its assessment of the likelihood that its kill rate is in fact at least 80%. Clearly, performing careful tests could allow the firm to learn whether the claim was true, allowing the firm to make that claim with confidence. Or, the firm could make such a claim on a speculative basis, not sure whether testing would bear out this assertion or not. Thus, the distinction between certainty and uncertainty is not central to the application of the model; what is central is the distinction between testable and untestable or verifiable and unverifiable to the regulator.

This example does, however, point to a feature of quality claims that is left outside of the present model, in which the standards of high and low quality are exogenous. In practice, the firm has some flexibility in defining the quality claim to be made, and this affects both how the consumers value the

high quality product and the exposure of the firm to penalties. To continue the herbicide example, the firm could instead make the claim that its product kills 60% of weeds. While this would probably reduce consumers' valuation of the product relative to the 80% claim (depending on the inferences they make based on the relative credibility of the two claims), the claim is also more likely to be true, which lowers the expected false claims penalty for a firm that was unsure of its product quality. Thus, the endogenous determination of the strength of the product quality claim could be an interesting feature to explore further in a similar model.

This analysis leaves aside many other important considerations by abstracting from the enforcement regime, incentives to innovate, the process of consumer expectations formation, and many other considerations that could be considered in future research. For example, the enforcement of these penalties is exogenous here; the policy-maker simply sets expected penalties for infringement of each prohibition. In fact, levels of enforcement and the endogenously determined expected penalties could hinge critically on the nature of the policy employed. If, for example, the substantiation requirement lowered government investigation and trial costs by shifting the burden of proof toward the firm compared to a scenario with only a false claims ban, then more stringent enforcement could be expected, meaning that a given nominal penalty might translate into a much higher expected penalty under that policy. Consideration of incentives to innovate would generally raise the social benefit to learning and dissemination of information, making policies that encourage learning and strong separation even more desirable. Similarly, one might doubt that weak separation equilibria are as desirable or as likely in real life as in the model, as they require sophisticated consumers to properly discount claims of product quality to sustain the equilibrium. In the absence of weak signaling as an alternative, learning and strong signaling become even more socially desirable, increasing the importance of the substantiation requirement.

7 References

- Akerlof, G., 1970, "The Market for 'Lemons': Quality Uncertainty and the Market Mechanism," *Quarterly Journal of Economics*, 488-500.
- Daughety, A. and J. Reinganum, 2008, "Products Liability, Signaling and Disclosure," *Journal of Institutional and Theoretical Economics*, 106-126.
- Farrell, J., 1986, "Voluntary Disclosure: Robustness of the Unraveling Result, and Comments on Its Importance," *Antitrust and Regulation*, 91-103.
- Grossman, S., 1981, "The Informational Role of Warranties and Private Disclosure about Product Quality," *Journal of Law and Economics*, 461-483.
- Matthews, S. and A. Postlewaite, 1985, "Quality Testing and Disclosure," *Rand Journal of Economics*, 328-340.
- Polinsky, M. and S. Shavell, 2006, "Mandatory Versus Voluntary Disclosure of Product Risks," NBER Working Paper 12776.
- Viscusi, K., 1978, "A Note on Lemons Markets with Quality Certification," *Rand Journal of Economics*, 277-279.